



مجموعه مقالات

هشتمین سمینار

آمار و احتمال فازی

«بزرگداشت استاد فقید پروفسور ناصر رضا ارقامی»

دانشکده علوم ریاضی دانشگاه فردوسی مشهد

۱۹ و ۲۰ اردیبهشت ماه ۱۳۹۷

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ

مجموعه مقالات

هشتمین سمینار

آمار و احتمال فازی

گروه آمار،

دانشکده علوم ریاضی،

دانشگاه فردوسی مشهد،

۱۹ و ۲۰ اردیبهشت ماه ۱۳۹۷

فهرست

۱.....	پیام دبیر هشتمین سمینار آمار و احتمال فازی
۲.....	چکیده ای از زندگینامه علمی پروفیسور ناصر رضا ارقامی
۳.....	محورهای سمینار و اعضای کمیته برگزارکننده سمینار
۴.....	اعضای کمیته علمی و مشاوران علمی سمینار
۵.....	کمیته اجرایی و کادر اجرایی سمینار
۶.....	حامیان سمینار

مقالات فارسی (به ترتیب الفبا بر اساس نام فامیل نویسنده اول)

۷.....	طرح نمونه گیری برای پذیرش دنباله ای برای مشخصه طول عمر در محیط فازی اصغری، ف. و صادقیورگیلده، ب.
۱۹.....	طرح اصلاحی نمونه گیری یک مرحله ای متغیر فازی افشاری، ر. و صادقیورگیلده، ب.
۲۵.....	کاربرد رگرسیون لجستیک فازی در مدل سازی شدت اختلال اوتیسم بهنامپور، ع.، بیگلریان، الف.، بخشی، ع. ا.، قربانی، م. و نامداری، م.
۳۴.....	برآورد پارامتر توزیع نمایی به کمک الگوریتم EM مبتنی بر مشاهدات فازی پرچمی، ع.
۳۹.....	پیشنهادی برای رفع چالش متغیر تصادفی فازی پرچمی، ع.
۴۷.....	رویکردی جدید به منظور مدلسازی و حل مسأله معکوس ۱-میانه فازی خداقلی، م.، دولتی، الف. و حسین زاده، ع.
۵۴.....	رگرسیون خطی فازی استوار براساس فاصله علامتدار بین دو عدد فازی خَمَر، الف. ح. و عارفی، م.
۶۲.....	نمودارهای کنترل فازی براساس پروفایل ها جلیلیان عاملی، م.، صادقیورگیلده، ب. و عباسی گنجی، ز.
۷۰.....	نمودار $RA-CUSUM$ بر اساس داده های فازی رضایی فر، الف.، صادقیورگیلده، ب. و محتشمی بزراداران، غ. ر.
۷۷.....	ضریب همبستگی در شرایط نادقیق زارعی، ر. و کیانپور، ط.
۸۸.....	مشارکت های استاد ارقامی در زمینه آمار و احتمال فازی طاهری، س. م.
۹۳.....	فازی سازی داده های باستان شناسی به منظور استفاده در آمار فازی طاهری، س. م.، ایروانی قدیم، ف. و کبیریان، م.
۹۹.....	رگرسیون خطی ساده چندکی فازی

عارفی، م.

- ۱۰۵..... دسترس پذیری سیستم های مهندسی با طول عمر فازی
عباس نیا، ز. ، توکلی ناریند، خ. و صادقپور گیلده، ب.
- ۱۱۲..... تصمیم گیری متعادل در آزمون فرضیه های فازی
مدنی، م. س.، ربیعی، م. ر. و پرچمی، ع.
- ۱۱۷..... غالب بودن شانس: معیاری برای مقایسه سیستم های پیچیده در فضای شانس
مهرعلیزاده، ر. الف. ، امینی، م. ، صادقپور گیلده، ب. و احمدزاده، ح.

مقالات انگلیسی (به ترتیب الفبا بر اساس نام فامیل نویسنده اول)

- Fuzzy process capability indices for simple linear profile** ۱۲۱
Abbasi Ganji, Z. and Sadeghpour Gildeh, B.
- A recurrent neural network model to solve the fuzzy shortest path problem** ۱۲۸
Eshaghnezhad, M., Rahbarnia, F. and Effati, S.
- JB estimation approach applied to fuzzy regression modeling** ۱۳۵
Kashani, M., Arashi, M. and Rabiei, M.R.
- Topic modeling based plagiarism dethection in Persian documents** ۱۴۲
Mohammadian, M., Khakpour, J. and Pedram, M. M.
- A capable neural network model for fuzzy quadratic optimization problems** ۱۴۹
Mansoori, A. and Effati, S.
- Fuzzy ridge linear regression analysis for fuzzy input-output data** ۱۵۶
Rabiei, M. R., Arashi, M., Goodarzi, V. and Piyadeh Kouhsar, F.
- A fuzzy optimal planned replacement time based on availability and cost functions** ۱۶۵
Safaei, F., Ahmadi, J. and Sadeghpour Gildeh, B.

پیام دبیر هشتمین سمینار تخصصی آمار و احتمال فازی

ارزش و اعتبار عالم و دانشمند به علم و دانش اوست و این یک تصادف نیست که در کتاب وحی الهی قرآن مجید، متجاوز از هفتصد مورد ماده علم و مشتقات آن به کار رفته و در اولین آیات نازل شده به رسول اکرم (ص) سخن از علم، دانش و قلم آمده است.

« همان کسی که به وسیله قلم تعلیم انسان نمود (سوره علق/آیه ۴) »

بهترین و زیباترین تعبیرها و تعریف‌ها درباره علم و دانش، در معارف دین ما ذکر شده و تصور نمی‌شود که در هیچ یک از مکاتب دیگر به اندازه مکتب اسلام به ارزش علم و عالم توجه شده باشد. علم وسیله‌ای است که در سایه آن اطاعت خدا محقق می‌شود و توحید در عبادت الهی به وسیله علم تحقق پیدا می‌کند و در یک کلمه خیر دنیا و آخرت در سایه علم تحصیل می‌شود و شر دنیا و آخرت به واسطه جهل و نادانی تحقق پیدا می‌کند. امام سجاد زین العابدین (ع) می‌فرماید: «حق معلم و آموزگار تو این است که او را بزرگ داری و مجلس او را محترم شماری و گفته‌های او را خوب گوش کنی و فراگیری، و او را کمک کنی تا آن چه را که از دانش و علم بدان نیازمندی به تو بیاموزد».

خداوند بزرگ را سپاسگزاریم که به ما نعمت خدمتگزاری به جمعی از اساتید و دانشجویان ایران عزیز و به ویژه مشارکت در مراسم بزرگداشت استاد فقید دکتر ناصررضا ارقامی، که از اسطوره‌های جامعه آماری (در زمینه‌های مختلف آمار از جمله آمار و احتمال فازی) بودند، را ارزانی نمود. جامعه آماردانان کشور شاهد یک عمر خدمت و تلاش صادقانه این استاد فرهیخته بود. بدون تردید همه دانش‌آموختگان آمار ایران در چند دهه گذشته در هر کسوتی به ویژه استادی و مدرسی دانشگاه‌های داخل و خارج مستقیم و غیر مستقیم از محضر این استاد فرزانه خوشه‌چینی نموده‌اند.

برای هر چه باشکوه‌تر برگزارکردن مراسم بزرگداشت زنده‌یاد دکتر ارقامی، با اساتید برجسته آمار کشور مکاتبات لازم برای شرکت در مراسم و ارسال پیام‌های متنی، صوتی یا تصویری انجام شد. از همه عزیزانی که با ارسال پیام، توجه خاص خود را نسبت به این استاد فقید ابراز داشتند، به ویژه از خانواده محترم استاد فقید شادوران دکتر ارقامی کمال تشکر را دارم.

استاد عزیز:

کاشکی می‌شد به جای زنده‌یاد می‌توانستیم بگوییم زنده باد

«روحش شاد و روانش غرق انوار الهی»

همت والای پژوهشگران حوزه آمار و احتمال فازی کشور با ارائه مقاله‌های وزین از یک طرف و تلاش همکاران و دانشجویان متعهد و دلسوز گروه آمار دانشگاه فردوسی از طرف دیگر موجبات اصلی برگزاری هشتمین سمینار آمار و احتمال فازی، در جوار بارگاه ملکوتی هشتمین اختر تابناک آسمان امامت و ولایت حضرت علی ابن موسی الرضا (ع)، و در دانشکده محل کار مرحوم استاد ارقامی، است. از اساتید ارجمند و دانشجویان گرامی که با مشارکت علمی فعال در این سمینار به بهتر برگزار شدن آن کمک کردند، سپاسگزارم. در خاتمه شایسته است مراتب قدردانی خود را از هیات رئیسه محترم دانشگاه فردوسی مشهد، هیات رئیسه محترم دانشکده علوم ریاضی دانشگاه فردوسی مشهد، هیأت مدیره محترم انجمن سیستم‌های فازی ایران، هیأت مدیره محترم انجمن آمار ایران، رئیس محترم مرکز مفاخر دانشگاه فردوسی مشهد، رئیس محترم دانشگاه پیام نور خراسان رضوی، هیات رئیسه محترم قطب علمی داده‌های ترتیبی و فضایی دانشگاه فردوسی مشهد، اعضای محترم هیات علمی، کارکنان عزیز و دانشجویان گرامی که با همت آن‌ها این همایش علمی به سرانجام رسید، اعلام نمایم.

بهرام صادقپور گیلده

دبیر هشتمین سمینار آمار و احتمال فازی

مشهد مقدس، اردیبهشت ۱۳۹۷

چکیده‌ای از زندگی‌نامه علمی پروفیسور ناصررضا ارقامی

دکتر ناصررضا ارقامی در سال ۱۳۲۸ در شهر نیشابور متولد شدند. دوران ابتدایی و متوسطه را در شهرستان نیشابور سپری کردند. پس از کسب مدرک دیپلم با استفاده از بورس بانک مرکزی در سال ۱۳۴۶ برای ادامه تحصیل به انگلستان سفر کردند و پس از اخذ مدرک کارشناسی در اقتصاد و کارشناسی ارشد در آمار سال ۱۳۵۰ به ایران بازگشتند و به استخدام دانشگاه فردوسی مشهد درآمدند. سال ۱۳۵۶ برای ادامه تحصیل در دوره دکتری به دانشگاه ایالتی فلوریدا آمریکا عازم شدند و پس از دفاع از رساله خود در سال ۱۳۶۰، به ایران بازگشتند. در سال ۱۳۸۸ پس از ۳۷ سال خدمت در دانشگاه فردوسی مشهد به درخواست خودشان بازنشسته شدند. زنده‌یاد در طول حیات با برکت خود سهم ارزنده‌ای در ارتقای جایگاه علم آمار کشور ایفا کردند. فراهم نمودن زمینه برای تحقیق و پژوهش جوانان در دوره دکتری در داخل کشور، تقویت بنیه علمی گروه آموزشی آمار، گسترش روابط بین المللی، ارتقای گروه آموزشی آمار به عنوان قطب علمی آمار کشور، تلاش برای اشاعه علم آمار از منظر کاربردی، راه‌اندازی مرکز مشاوره آماری، برگزاری کارگاه‌های آموزشی کاربردی، برگزاری همایش بین المللی داده‌های ترتیبی در دانشگاه فردوسی مشهد (۲۰۰۶)، تألیف و ترجمه کتاب‌های آماری، چاپ و نشر بیش از ۱۰۰ مقاله علمی پژوهشی، راهنمایی حدود ۴۰ پایان‌نامه کارشناسی ارشد و ۱۲ رساله دکتری آمار، بخشی از خدمات ارزنده این استاد فقید است.

آقای دکتر ناصررضا ارقامی ظهر روز چهارشنبه ۵ مرداد ۱۳۹۶ در محل دفتر کار خود در دانشکده علوم ریاضی، بر اثر سکته قلبی دارفانی را وداع نمود و کشور را داغدار کرد.

یادش گرامی و روحش شاد باد.

محورهای سمینار

- احتمال فازی
- رگرسیون فازی
- بهینه سازی فازی
- کنترل کیفیت فازی
- قابلیت اعتماد در محیط فازی
- استنباط آماری در محیط فازی
- وجوه آماری و احتمالی سیستمهای فازی
- فرایندهای تصادفی و سریهای زمانی در محیط فازی
- آمار حیاتی و آمار اقتصادی-اجتماعی در محیط فازی
- سایر موضوعهای مرتبط با آمار و احتمال در محیط فازی

اعضای کمیته برگزارکننده

- دکتر علیرضا سهیلی، دانشگاه فردوسی مشهد (رئیس دانشکده علوم ریاضی)
- دکتر سید محمود طاهری، دانشگاه تهران
- دکتر جعفر احمدی، دانشگاه فردوسی مشهد
- دکتر محمد امینی، دانشگاه فردوسی مشهد
- دکتر مهدی جباری نوقابی، دانشگاه فردوسی مشهد
- دکتر آرزو حبیبی راد، دانشگاه فردوسی مشهد
- دکتر مهدی دوست پرست، دانشگاه فردوسی مشهد
- دکتر مصطفی رزمخواه، دانشگاه فردوسی مشهد
- دکتر عبدالحمید رضائی رکن آبادی، دانشگاه فردوسی مشهد
- دکتر مجید سرمد، دانشگاه فردوسی مشهد
- دکتر بهرام صادقپور گیلده، دانشگاه فردوسی مشهد (دبیر کمیته برگزاری سمینار)
- دکتر مهدی عمادی، دانشگاه فردوسی مشهد

- دکتر معصومه فشندی، دانشگاه فردوسی مشهد
- دکتر وحید فکور، دانشگاه فردوسی مشهد
- دکتر غلامرضا محتشمی برزادران، دانشگاه فردوسی مشهد

اعضای کمیته علمی (به ترتیب حروف الفبا)

- دکتر جعفر احمدی، دانشگاه فردوسی مشهد
- دکتر محمدرضا اکبرزاده توتونچی، دانشگاه فردوسی مشهد
- دکتر محمد امینی، دانشگاه فردوسی مشهد
- دکتر محمد رضا ربیعی، دانشگاه صنعتی شاهرود
- دکتر عبدالحمید رضائی رکنآبادی، دانشگاه فردوسی مشهد
- دکتر بهرام صادقیپور گیلده، دانشگاه فردوسی مشهد (دبیر کمیته علمی سمینار)
- دکتر سید محمود طاهری، دانشگاه تهران
- دکتر محسن عارفی، دانشگاه بیرجند
- دکتر بهروز فتحی واجارگاه، دانشگاه گیلان
- دکتر ماشاالله ماشینچی، دانشگاه شهید باهنر کرمان
- دکتر غلامرضا محتشمی برزادران، دانشگاه فردوسی مشهد
- دکتر علی وحیدیان کامیاد، دانشگاه فردوسی مشهد

اعضای کمیته مشاوران علمی سمینار (به ترتیب حروف الفبا)

- دکتر حامد احمدزاده، دانشگاه سیستان و بلوچستان
- دکتر محمدقاسم اکبری، دانشگاه بیرجند
- دکتر میرمحسن پدرام، دانشگاه خوارزمی
- دکتر عباس پرچی، دانشگاه شهید باهنر کرمان
- دکتر جلال چاچی، دانشگاه سمنان
- دکتر رضا زارعی، دانشگاه گیلان

- دکتر زینب عباسی گنجی، مرکز تحقیقات جهاد کشاورزی مشهد
- دکتر سید باقر میراشرفی لنگرودی، دانشگاه مازندران
- دکتر حمیدرضا نیلی ثانی، دانشگاه بیرجند

اعضای کمیته اجرایی

- دکتر جعفر احمدی، دانشگاه فردوسی مشهد
- دکتر محمد امینی، دانشگاه فردوسی مشهد
- دکتر هادی جباری نوقابی، دانشگاه فردوسی مشهد
- دکتر مهدی جباری نوقابی، دانشگاه فردوسی مشهد
- دکتر مجید سرمد، دانشگاه فردوسی مشهد
- دکتر بهرام صادقپور گیلده، دانشگاه فردوسی مشهد (دبیر کمیته اجرایی سمینار)
- دکتر غلامرضا محتشمی برزادران، دانشگاه فردوسی مشهد

کادر اجرایی سمینار

- حسینعلی محتشمی برزادران
- عادل احمدی نادی
- مهدیه عرفانیان
- فاطمه صفائی
- مرتضی محمدی
- منصوره روشن نژاد
- مجتبی اصفهانی
- جابر کاظم پور
- حسن حیدری

حامیان سمینار آمار و احتمال فازی

- انجمن سیستم‌های فازی ایران
- انجمن آمار ایران
- قطب علمی داده‌های ترتیبی و فضایی
- دانشگاه پیام نور خراسان رضوی
- پایگاه استنادی علوم جهان اسلام
- مرکز آثار مفاخر و اسناد دانشگاه فردوسی مشهد
- شرکت شهرک‌های صنعتی خراسان رضوی

طرح نمونه گیری برای پذیرش دنباله‌ای برای مشخصه طول عمر در محیط فازی

سیده فاطمه اصغری بایگی، بهرام صادقیپور گیلده

دانشگاه فردوسی مشهد، گروه آمار

چکیده

در این مقاله، یک طرح نمونه گیری جدید بر مبنای آماره نسبت احتمال دنباله‌ای برای آزمون فرضیه‌های فازی معرفی می‌کنیم. برای مقابله با عدم حتمیت رخ داده، از اعداد فازی مثلثی (TFN) برای بیان پارامترهای فازی، یعنی سطح کیفی قابل قبول (AQL) و درصد معیوب‌های محموله ($LTPD$)، استفاده می‌شود. معیار تصمیم‌گیری برای پذیرش و یا رد را برای λ -برش دلخواه طراحی کرده ایم.

کلمات کلیدی:

۱ مقدمه

طرح نمونه‌گیری پذیرشی محموله ($LASP$)^۱، یکی از ابزارهای مهم کنترل کیفیت آماری (SQC)^۲ است که می‌تواند به روش‌های مختلفی مورد استفاده قرار گیرد. در نمونه برداری دنباله‌ای، هر بار نمونه انتخاب می‌شود و با استفاده از آن و نمونه‌های قبلی آماره آزمون محاسبه می‌گردد، این کار تا زمانی که تصمیم به رد یا پذیرش محموله اتخاذ شود، ادامه پیدا می‌کند.

اگر اندازه نمونه یک باشد، آن را نمونه‌گیری مورد به مورد (یکی یکی) دنباله‌ای^۳ می‌نامند. در غیر این صورت، یعنی زمانی که اندازه نمونه (در هر بار نمونه برداری) بیشتر از یک باشد، طرح نمونه برداری دنباله‌ای گروهی نامیده می‌شود. نمونه‌گیری دنباله‌ای در سال ۱۹۴۳ برای استفاده در بازرسی کیفی سریع در صنایع دفاعی توسعه یافت [۱۶]. این طرح توسط بسیاری از محققان از جمله هاردو سحایی^۴ و همکارانش [۲۱] مورد مطالعه قرار گرفت. طرح نمونه‌گیری دنباله‌ای یکی یکی (مورد به مورد) در فرم سنتی براساس آزمون نسبت احتمال دنباله‌ای ($SPRT$)^۵ با فرضیه‌های دقیق طراحی می‌شود (برای جزئیات بیشتر به [۲] مراجعه شود). با این حال، موقعیت‌های مختلفی وجود دارد که پارامترهای طرح دقیقاً مشخص نیستند. در طراحی طرح دنباله‌ای دو سطح کیفیت AQL و $LTPD$ معلوم هستند. از سوی دیگر، گاهی اوقات پارامتر کیفیت در دسترس دقیق نیست. برای غلبه بر این مشکل ابتدا $SPRT$ را برای فرضیه‌های فازی در حالتی که پارامترهای نامعلوم آنها به صورت عدهای فازی

^۱plan sampling acceptance lot the

^۲control quality statistical

^۳sampling sequential item by item

^۴Hardeo

^۵test ratio probability sequential the

هستند، توسعه می دهیم. سپس نمونه گیری برای پذیرش مورد به مورد دنباله ای تعریف شده براساس SPRT فازی طراحی می کنیم. آزمون فرضیه های فازی اولین بار توسط آرنولد، دلگادو^۶ و همکارانش مطرح شد. واتانبه و امیزومی^۷ [۱۹]، طاهری و بهبودیان [۲۰]، ترابی و همکاران [۱۱]، ترابی و بهبودیان [۱۲]، ترابی و میر حسینی [۱۳]، ترابی و بهبودیان [۱۰]، اکبری [۱۷]، صادق پورگیلده و همکاران [۵] و جمخانه بالویی و همکاران [۹-۷] روش های آزمون فرضیه های فازی را در طراحی نمونه گیری برای پذیرش استفاده کردند. ساختار مقاله به صورتی است که در بخش بعد آزمون نسبت احتمال دنباله ای برای فرضیه های فازی را یادآوری می کند. در بخش سوم، طرح نمونه گیری برای پذیرش دنباله ای فازی را طراحی کرده و خطوط محدود کننده (تصمیم گیری) آن را محاسبه می کنیم. در ادامه بخش سوم در مورد داده های طول عمر شبیه سازی انجام شده است.

۲ SPRT آزمون نسبت احتمال دنباله ای برای فرضیه های فازی

در آزمون فرضیه های فازی، فرضیه ها به صورت زیراند:

$$\begin{cases} H^{\circ} : \theta \text{ is } H^{\circ}(\theta), \\ H^{\wedge} : \theta \text{ is } H^{\wedge}(\theta), \end{cases} \quad (۱)$$

یا

$$\begin{cases} H^{\circ} : \tilde{\theta} = \tilde{\theta}^{\circ}, \\ H^{\wedge} : \tilde{\theta} = \tilde{\theta}^{\wedge} \end{cases} \quad (۲)$$

که در آن $j = 0, 1$ ، $H_j(\theta)$ اشاره دارد که θ یک زیر مجموعه فازی از Θ (فضای پارامتر) با تابع عضویت $H_j(\theta)$ است و می نویسیم $H_j(\theta) : \Theta \rightarrow [0, 1]$.

فرض کنید متغیر تصادفی X دارای تابع چگالی احتمال فازی (FPDF)^۸، $\tilde{f}(x, \tilde{\theta})$ که $\tilde{\theta}$ یک پارامتر فازی است، باشد. تحت فرضیه $j = 0, 1$ ، $H_j(\theta)$ ، λ -برش برای چگالی احتمال فازی برابر است با:

$$\tilde{f}_j(x, \theta_j) [\lambda] = \{f_j(x, \theta_j) | \theta_j \in \tilde{\theta}_j [\lambda]\} = \{f_j^L [\lambda], f_j^U [\lambda]\} \quad j = 0, 1, \quad (۳)$$

که در آن

$$\begin{aligned} f_j^L(x) [\lambda] &= \min\{f_j(x, \theta_j) | \theta_j \in \tilde{\theta}_j [\lambda]\}, \\ f_j^U(x) [\lambda] &= \max\{f_j(x, \theta_j) | \theta_j \in \tilde{\theta}_j [\lambda]\}. \end{aligned} \quad (۴)$$

حال فرض کنید $X = (X_1, X_2, \dots, X_n)$ نمونه ای تصادفی از یک جامعه با تابع چگالی احتمال فازی $\tilde{f}(x; \tilde{\theta})$ و تابع آزمون $\Phi(X)$ باشد. احتمال خطای نوع اول و دوم برای آزمون فرض فازی (۲) به صورت زیر است:

$$\tilde{\beta}_{\Phi} = 1 - \tilde{E}^{\wedge}_1(\Phi(X)) \text{ و } \tilde{\alpha}_{\Phi} = \tilde{E}^{\circ}(\Phi(X))$$

^۶Arnold and Delgado

^۷Watanabe and Imaizumi

^۸function density probabiliti fuzzy the

که در آن $\tilde{E}_j(\Phi(X))$ مقدار مورد انتظار فازی تابع آزمون تحت تابع چگالی، $j = 0, 1$ ، $\tilde{f}_j(x, \tilde{\theta})$ است.

$$\begin{aligned}\tilde{E}_0(\Phi(x))[\lambda] &= \left\{ \int_x \Phi(x) f_j(x, \theta_j) dx | \theta_j \in \tilde{\theta}_j[\lambda] \right\} \\ &= [E_j^L(\Phi(X))[\lambda], E_j^U(\Phi(X))[\lambda]], j = 0, 1,\end{aligned}\quad (5)$$

که

$$\begin{aligned}E_j^L &= \min \left\{ \int_x \Phi(x) f_j(x, \theta_j) dx | \theta_j \in \tilde{\theta}_j[\lambda] \right\}, \\ E_j^U &= \max \left\{ \int_x \Phi(x) f_j(x, \theta_j) dx | \theta_j \in \tilde{\theta}_j[\lambda] \right\}.\end{aligned}\quad (6)$$

حال آزمون نسبت احتمال دنباله ای برای FHT⁹ را پیشنهاد می کنیم. فرض کنید $X_1, X_2, \dots$ یک دنباله از متغیرهای تصادفی iid از جامعه ای با تابع چگالی احتمال فازی $\tilde{f}(x, \tilde{\theta})$ باشد، در این صورت λ -برش نسبت احتمال دنباله ای با استفاده از روش بالکلی به صورت زیر است :

$$\begin{aligned}\tilde{R}_n(x) &= \left\{ \frac{L_{\tilde{\theta}_0}(x_1, \dots, x_n)}{L_{\tilde{\theta}_1}(x_1, \dots, x_n)} | \theta_j \in \tilde{\theta}_j[\lambda], j = 0, 1 \right\}, \\ &= \left\{ \frac{\tilde{f}_0(x_i, \tilde{\theta}_0)}{\tilde{f}_1(x_i, \tilde{\theta}_1)} | \theta_j \in \tilde{\theta}_j[\lambda], j = 0, 1 \right\}, \\ &= \left[\frac{\prod_{i=1}^n f_0^L(x_i, \theta_0^L)}{\prod_{i=1}^n f_1^U(x_i, \theta_1^U)}, \frac{\prod_{i=1}^n f_0^U(x_i, \theta_0^U)}{\prod_{i=1}^n f_1^L(x_i, \theta_1^L)} \right], \\ &= \left[\frac{L_{\tilde{\theta}_0}^L}{L_{\tilde{\theta}_1}^U}, \frac{L_{\tilde{\theta}_0}^U}{L_{\tilde{\theta}_1}^L} \right], \\ &= [R_n^L(x), R_n^U(x)].\end{aligned}\quad (7)$$

که

$$\begin{aligned}f_j^L(x_i, \theta_j^L) &= \min \left\{ f_j(x_i, \theta_j) | \theta_j \in \tilde{\theta}_j[\lambda] \right\}, j = 0, 1, \\ f_j^U(x_i, \theta_j^U) &= \max \left\{ f_j(x_i, \theta_j) | \theta_j \in \tilde{\theta}_j[\lambda] \right\}, j = 0, 1.\end{aligned}\quad (8)$$

الگوریتم آزمون: برای k_0 و k_1 ثابت ($0 < k_0 < k_1$) انجام آزمون به صورت زیر است :

مشاهده x_1 گرفته و $\tilde{R}_1$ محاسبه شود؛

اگر $R_1^U \leq k_0$ آنگاه H_0 را رد می کنیم،

اگر $R_1^L \geq k_1$ آنگاه H_0 را می پذیریم،

و اگر $k_1 < R_1^L$ و $R_1^U > k_0$ مشاهده x_2 را انتخاب کرده و $\tilde{R}_2$ را محاسبه می کنیم؛

اگر $R_2^U \leq k_0$ آنگاه H_0 را رد می کنیم،

اگر $R_2^L \geq k_1$ آنگاه H_0 را می پذیریم،

و اگر $k_1 < R_2^L$ و $R_2^U > k_0$ مشاهده x_3 را می گیریم؛

و نمونه گیری را تا زمانی که $k_1 < R_n^L$ و $R_n^U > k_0$ ادامه می دهیم و به محض اینکه $R_n^U \leq k_0$ یا $R_n^L \geq k_1$

متوقف می شویم.

⁹testing hypothesis fuzzy

H_0 را رد می‌کنیم اگر $R_n^U \leq k_0$ و می‌پذیریم اگر $R_n^U > k_0$.
در نتیجه ناحیه بحرانی یا رد فرضیه H_0 ، برای آزمون فرضیه آماری به صورت $C = \cup_{n=1}^{\infty} C_n$ که،

$$C_n = \{(x_1, \dots, x_n) | R_j^L < k_1 \ \& \ R_j^U > k_0, j = 1, \dots, n-1, R_n^U \leq k_0\}. \quad (9)$$

به طور مشابه ناحیه پذیرش می‌تواند به صورت $A = \cup_{n=1}^{\infty} A_n$ ، که در آن

$$A_n = \{(x_1, \dots, x_n) | R_j^L < k_1 \ \& \ R_j^U > k_0, j = 1, \dots, n-1, R_n^L \geq k_1\}. \quad (10)$$

بنابراین λ -برش اندازه‌ی خطاهای نوع اول و دوم این آزمون برابر است با:

$$\begin{aligned} \tilde{\alpha}[\lambda] &= \left\{ \sum_{n=1}^{\infty} \int_{C_n} L_{\tilde{\theta}_0}(x) dx | \theta_0 \in \tilde{\theta}_0[\lambda] \right\}, \\ \tilde{\beta}[\lambda] &= \left\{ \sum_{n=1}^{\infty} \int_{A_n} L_{\tilde{\theta}_1}(x) dx | \theta_1 \in \tilde{\theta}_1[\lambda] \right\}. \end{aligned} \quad (11)$$

نکته: فرض کنید

$$\tilde{A}[\lambda] = [A^L[\lambda], A^U[\lambda]], \quad \tilde{B}[\lambda] = [B^L[\lambda], B^U[\lambda]]$$

λ -برش دو عدد فازی $\tilde{A}$ و $\tilde{B}$ باشد گوییم $\tilde{A} < \tilde{B}$ اگر و فقط اگر

$$\forall \lambda \in [0, 1], \quad A^U[\lambda] < B^L[\lambda] \text{ باشد.}$$

لم ۱.۲. فرض کنید k_0 و k_1 به گونه‌ای تعریف شوند که $SPRT$ برای فرضیه فازی دارای خطای نوع اول و دوم ثابت $\tilde{\alpha}$ و $\tilde{\beta}$ به ترتیب باشد، سپس k_0 و k_1 میتواند به وسیله

$$k'_0 = \frac{\alpha^U}{1 - \beta^U} \quad ; \quad k'_1 = \frac{1 - \alpha^U}{\beta^U} \quad (12)$$

تقریب زده شود.

لم ۲.۲. اگر $\tilde{\alpha}$ و $\tilde{\beta}$ به ترتیب، اندازه خطای نوع اول و دوم $SPRT$ در صورت استفاده از k'_0 و k'_1 باشند، آنگاه برای هر λ -برش دلخواه داریم:

$$\alpha'^U + \beta'^U \leq \alpha^U + \beta^U.$$

برای مشاهده اثبات لم‌های بالا به [۱] مراجعه نمایید. اگر

$$\begin{aligned} \tilde{Z}_i &= [Z_i^L, Z_i^U], \\ &= \left[\ln \frac{f_0^L(x_i, \theta_0^L)}{f_1^U(x_i, \theta_1^U)}, \ln \frac{f_0^U(x_i, \theta_0^U)}{f_1^L(x_i, \theta_1^L)} \right]. \end{aligned} \quad (13)$$

با مشاهده‌ی یک نمونه و استفاده از رابطه‌ی (۱۳)، یک آزمون معادل برای آزمون نسبت احتمال دنباله‌ای فازی ارایه می‌شود:

نمونه‌گیری را تا زمانی که $\sum_{i=1}^n Z_i^L < \ln(k'_1)$ و $\ln(k'_0) < \sum_{i=1}^n Z_i^U$ ادامه می‌دهیم، در غیر این صورت نمونه‌گیری متوقف می‌شود.

اگر $\sum_{i=1}^n Z_i^L \geq \ln(k'_1)$ ، H_0 را می‌پذیریم و اگر $\sum_{i=1}^n Z_i^U \leq \ln(k'_0)$ ، H_0 رد خواهد شد.

نمونه‌گیری را پایان می‌دهیم و H_0 را رد کرده یا می‌پذیریم، هرگاه برای اولین n داشته باشیم $R_n^L \geq k_1$ یا $R_n^U \leq k_0$. البته از پیش نمی‌دانیم که اندازه نمونه لازم برای رد یا پذیرش H_0 چقدر است. اندازه این نمونه یک متغیر تصادفی N می‌باشد، به طوری که پیشامد $(N \leq n)$ به $X_1, \dots, X_n$ بستگی دارد ولی مستقل

از $X_{n+1}, X_{n+2}, \dots$ می باشد. مقادیر این متغیر تصادفی که پایان نمونه گیری را اعلان می دارد، می تواند $1, 2, \dots$ باشد و به متغیر ایست^{۱۰} (زمان توقف) معروف است. پس

$$\tilde{E}_0^{(l)} \leq \tilde{E}(N|H_0 \text{ is true}) \leq \tilde{E}_0^{(u)}. \quad (14)$$

که در آن

$$\begin{aligned} \tilde{E}_0^{(l)} &= \frac{\tilde{\alpha} L n \frac{\alpha^U}{1 - \beta^U} + (1 - \tilde{\alpha}) L n \frac{1 - \alpha^U}{\beta^U}}{\tilde{E}(Z_i^U | H_0)}, \\ \tilde{E}_0^{(u)} &= \frac{\tilde{\alpha} L n \frac{\alpha^U}{1 - \beta^U} + (1 - \tilde{\alpha}) L n \frac{1 - \alpha^U}{\beta^U}}{\tilde{E}(Z_i^L | H_0)}, \end{aligned} \quad (15)$$

و

$$\tilde{E}_1^{(l)} \leq \tilde{E}(N|H_1) \leq \tilde{E}_1^{(u)} \quad (16)$$

به طوری که

$$\begin{aligned} \tilde{E}_1^{(l)} &= \frac{(1 - \beta) L n \frac{\alpha^U}{1 - \beta^U} + \tilde{\beta} L n \frac{1 - \alpha^U}{\beta^U}}{\tilde{E}(Z_i^U | H_1)}, \\ \tilde{E}_1^{(u)} &= \frac{(1 - \beta) L n \frac{\alpha^U}{1 - \beta^U} + \tilde{\beta} L n \frac{1 - \alpha^U}{\beta^U}}{\tilde{E}(Z_i^L | H_1)}, \end{aligned} \quad (17)$$

و

$$\tilde{E}(Z_i^{(l)} | H_j \text{ is true}) = \left\{ \int_x L n \frac{f_0^L(x, \theta_0^L)}{f_1^U(x, \theta_1^U)} f_j(x, \theta_j) dx | \theta_j \in \tilde{\theta}_j[\lambda] \right\}. \quad (18)$$

$$\tilde{E}(Z_i^{(U)} | H_j \text{ is true}) = \left\{ \int_x L n \frac{f_0^U(x, \theta_0^L)}{f_1^L(x, \theta_1^U)} f_j(x, \theta_j) dx | \theta_j \in \tilde{\theta}_j[\lambda] \right\}. \quad (19)$$

۳ طرح نمونه گیری برای پذیرش دنباله ای در محیط فازی براساس داده های طول عمر

خطوط پذیرش و رد، معیار تصمیم گیری در یک نمونه گیری دنباله ای یکی یکی (مورد به مورد) است.

شکل ۱ نحوه ی اجرای چنین طرح نمونه گیری را نشان میدهد. روش طرح به شرح زیر است:

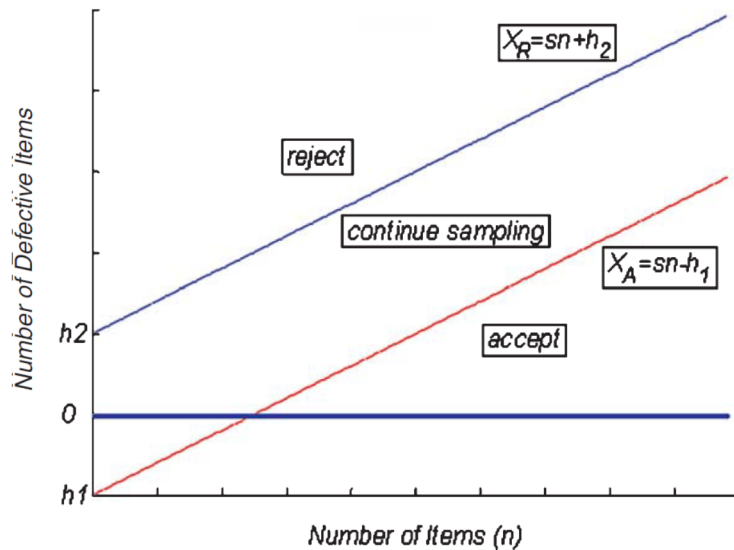
(۱) اگر نقطه ترسیم شده درون خطوط محدودیت قرار گیرد، آنگاه فرایند با رسم یک نمونه دیگر ادامه می یابد.

(۲) هنگامی که نقطه ترسیم شده بالای خط رد X_R قرار گیرد، H_0 رد می شود.

(۳) هنگامی که نقطه ترسیم شده زیر خط پذیرش X_A قرار گیرد، H_1 پذیرفته می شود.

فرضیه های H_0 و H_1 به صورت زیر بازنویسی می شوند:

^{۱۰} variable Stopping



شکل ۱: نمودار روند انجام طرح نمونه‌گیری دنباله‌ای

H_0 : سطح کیفیت قابل قبول انباشته (AQL)

(۲۰)

H_1 : سطح کیفیت قابل رد انباشته (RQL)

AQL حداکثر ارقام معیوبی است که جهت نمونه‌گیری برای پذیرش می‌تواند رضایت بخش باشد. نام جایگزین برای RQL درصد تحمل معیوب‌های محموله LTPD است.

۴ آزمون نسبت احتمال دنباله‌ای فازی بر اساس داده‌های طول عمر

توزیع‌های متعددی درباره‌ی پدیده‌های تصادفی مرتبط با قابلیت اعتماد ارایه شده است، که توزیع‌نمایی معروف‌ترین آنهاست. توزیع‌نمایی یکی از مهم‌ترین توزیع‌های آماری است که در مدل‌سازی و تحلیل داده‌های طول عمر به کار می‌رود. از جمله دلایل مهم کاربرد فراوان توزیع‌نمایی در قابلیت اعتماد، فرم ساده‌ی محاسباتی این توزیع و خواص منحصر به فرد این مدل است. متغیر تصادفی طول عمر T دارای توزیع‌نمایی با پارامتر مقیاس θ است، و به طور خلاصه می‌نویسیم $T \sim E(\theta)$ ، اگر تابع چگالی آن به صورت زیر باشد:

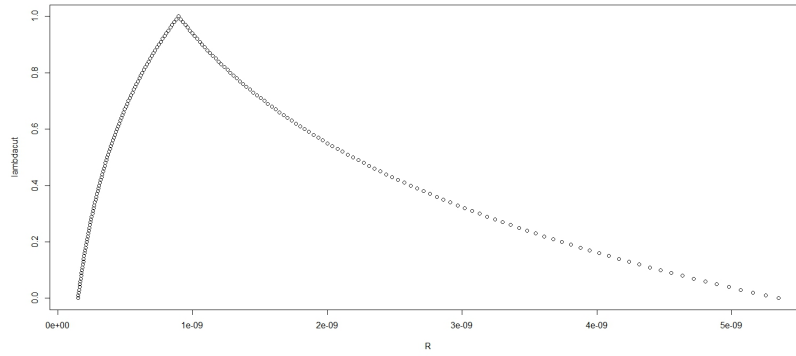
$$f(t, \theta) = \theta e^{-\theta t} \quad \theta > 0; t > 0. \quad (21)$$

حال فرضیه‌های فازی به ازای $j = 0, 1$ ، معادل (۳۰۱) را به صورت زیر بازنویسی می‌کنیم:

$$H_j(\theta) = \begin{cases} \frac{\theta - \theta_j^{(1)}}{\theta_j^{(2)} - \theta_j^{(1)}} & , \quad \theta_j^{(1)} \leq \theta \leq \theta_j^{(2)} \\ \frac{\theta_j^{(2)} - \theta}{\theta_j^{(2)} - \theta_j^{(1)}} & , \quad \theta_j^{(2)} \leq \theta \leq \theta_j^{(3)} \end{cases} \quad (22)$$

$$\forall \theta > 0, j = 0, 1.$$

در نهایت با استفاده از رابطه (۷۰۲) و (۲۰۴) برای هر λ -برش دلخواه، داریم:



شکل ۲: نمودار تابع عضویت $\tilde{R}_n$ در آزمون نسبت احتمال دنباله ای فازی

$$R_n^L(x) = \min \left\{ \frac{L_o(x, \theta_o)}{L_{\setminus}(x, \theta_{\setminus})} | \theta_o \in \tilde{\theta}_o[\lambda], \theta_{\setminus} \in \tilde{\theta}_{\setminus}[\lambda] \right\},$$

$$= \frac{\min \prod_{i=1}^n \left\{ \theta_o e^{-\theta_o x_i} | \theta_o \in \tilde{\theta}_o[\lambda] \right\}}{\max \prod_{i=1}^n \left\{ \theta_{\setminus} e^{-\theta_{\setminus} x_i} | \theta_{\setminus} \in \tilde{\theta}_{\setminus}[\lambda] \right\}} \quad (23)$$

و

$$R_n^U(x) = \max \left\{ \frac{L_o(x, \theta_o)}{L_{\setminus}(x, \theta_{\setminus})} | \theta_o \in \tilde{\theta}_o[\lambda], \theta_{\setminus} \in \tilde{\theta}_{\setminus}[\lambda] \right\},$$

$$= \frac{\max \prod_{i=1}^n \left\{ \theta_o e^{-\theta_o x_i} | \theta_o \in \tilde{\theta}_o[\lambda] \right\}}{\min \prod_{i=1}^n \left\{ \theta_{\setminus} e^{-\theta_{\setminus} x_i} | \theta_{\setminus} \in \tilde{\theta}_{\setminus}[\lambda] \right\}} \quad (24)$$

با توجه به رابطه های بالا چون نمی توان به طور تحلیلی مقادیر کمینه و بیشینه را پیدا کرد، لازم است محاسبات مربوطه را به کمک ماشین انجام دهیم. یعنی باید مقادیر R_n^U و R_n^L را به ازای چهار مقدار θ_o^L و θ_o^U و $\theta_{\setminus}^L$ و $\theta_{\setminus}^U$ به دست آوریم. فرضیه های زیر را در نظر میگیریم:

$$\begin{cases} H_o : \tilde{\theta}_o = 0.44, \\ H_{\setminus} : \tilde{\theta}_{\setminus} = 0.7, \end{cases} \quad (25)$$

به عبارت دیگر:

$$H_o(\theta) = \begin{cases} \frac{\theta - 0.44}{0.5 - 0.44} & , \quad 0.44 \leq \theta \leq 0.5 \\ \frac{0.56 - \theta}{0.56 - 0.5} & , \quad 0.5 \leq \theta \leq 0.56 \end{cases} \quad (26)$$

$$H_{\setminus}(\theta) = \begin{cases} \frac{\theta - 0.64}{0.7 - 0.64} & , \quad 0.64 \leq \theta \leq 0.7 \\ \frac{0.76 - \theta}{0.76 - 0.7} & , \quad 0.7 \leq \theta \leq 0.76 \end{cases} \quad (27)$$

شکل ۲ نمودار تابع عضویت $\tilde{R}_n$ در آزمون نسبت احتمال دنباله ای فازی را نشان می دهد.

n	λ -برش	$\tilde{R}_n(x) = [R_n^L(x), R_n^U(x)]$	تصمیم
1	0.0	[0.400513 , 0.5829364]	ادامه نمونه گیری
	0.1	[0.4074872 , 0.5713344]	ادامه نمونه گیری
	0.2	[0.4145773 , 0.5599459]	ادامه نمونه گیری
	0.3	[0.4217872 , 0.5487651]	ادامه نمونه گیری
	0.4	[0.4291196 , 0.5377864]	ادامه نمونه گیری
	0.5	[0.4365778 , 0.5270043]	ادامه نمونه گیری
	0.6	[0.4441647 , 0.5164137]	ادامه نمونه گیری
	0.7	[0.4518836 , 0.5060093]	ادامه نمونه گیری
	0.8	[0.4597379 , 0.4957864]	ادامه نمونه گیری
	0.9	[0.4677310 , 0.4857403]	ادامه نمونه گیری
	1.0	[0.4758665 , 0.4758665]	ادامه نمونه گیری
2	0.0	[0.06158364 , 0.13252215]	ادامه نمونه گیری
	0.1	[0.06384689 , 0.12729955]	ادامه نمونه گیری
	0.2	[0.06619185 , 0.12227517]	ادامه نمونه گیری
	0.3	[0.06862174 , 0.11744084]	ادامه نمونه گیری
	0.4	[0.07113992 , 0.11278875]	ادامه نمونه گیری
	0.5	[0.07374989 , 0.10831148]	ادامه نمونه گیری
	0.6	[0.07645533 , 0.10400197]	ادامه نمونه گیری
	0.7	[0.07926008 , 0.09985346]	ادامه نمونه گیری
	0.8	[0.08216815 , 0.09585956]	ادامه نمونه گیری
	0.9	[0.08518375 , 0.09201412]	ادامه نمونه گیری
	1.0	[0.0883113 , 0.0883113]	ادامه نمونه گیری
3	0.0	[0.02803753 , 0.08763024]	ادامه نمونه گیری
	0.1	[0.02956781 , 0.08250145]	ادامه نمونه گیری
	0.2	[0.03118059 , 0.07766560]	ادامه نمونه گیری
	0.3	[0.03288048 , 0.07310549]	ادامه نمونه گیری
	0.4	[0.03467237 , 0.06880499]	ادامه نمونه گیری
	0.5	[0.03656146 , 0.06474899]	فرض صفر رد می شود.
	0.6	[0.03855326 , 0.06092332]	فرض صفر رد می شود.
	0.7	[0.04065361 , 0.05731469]	فرض صفر رد می شود.
	0.8	[0.04286873 , 0.05391063]	فرض صفر رد می شود.
	0.9	[0.04520521 , 0.05069942]	فرض صفر رد می شود.
	1.0	[0.04767006 , 0.04767006]	فرض صفر رد می شود.

شکل ۳: تصمیم‌گیری براساس آزمون نسبت احتمال دنباله ای فازی

$$\tilde{\theta}_0 \approx 0.5, \tilde{\theta}_1 \approx 0.7, \tilde{\alpha} \approx 0.5, \tilde{\beta} \approx 0.1$$

4	0.0	[0.01108755 , 0.05044801]	فرض صفر رد می شود.
	0.1	[0.01189654 , 0.04655012]	فرض صفر رد می شود.
	0.2	[0.01276399 , 0.04294807]	فرض صفر رد می شود.
	0.3	[0.01369421 , 0.03961917]	فرض صفر رد می شود.
	0.4	[0.01469185 , 0.03654253]	فرض صفر رد می شود.
	0.5	[0.01576191 , 0.03369893]	فرض صفر رد می شود.
	0.6	[0.01690977 , 0.03107064]	فرض صفر رد می شود.
	0.7	[0.01814126 , 0.02864134]	فرض صفر رد می شود.
	0.8	[0.01946263 , 0.02639599]	فرض صفر رد می شود.
	0.9	[0.02088067 , 0.02432070]	فرض صفر رد می شود.
	1.0	[0.02240266 , 0.02240266]	فرض صفر رد می شود.
5	0.0	[0.003521026 , 0.023406337]	فرض صفر رد می شود.
	0.1	[0.003845172 , 0.021167980]	فرض صفر رد می شود.
	0.2	[0.004198927 , 0.019140708]	فرض صفر رد می شود.
	0.3	[0.004585029 , 0.017304544]	فرض صفر رد می شود.
	0.4	[0.005006472 , 0.015641440]	فرض صفر رد می شود.
	0.5	[0.00546654 , 0.01413509]	فرض صفر رد می شود.
	0.6	[0.005968829 , 0.012770742]	فرض صفر رد می شود.
	0.7	[0.006517282 , 0.011535067]	فرض صفر رد می شود.
	0.8	[0.007116226 , 0.010415996]	فرض صفر رد می شود.
	0.9	[0.007770403 , 0.009402610]	فرض صفر رد می شود.
	1.0	[0.008485023 , 0.008485023]	فرض صفر رد می شود.
...			
15	0.0	[1.378882e-13 4.985301e-11]	فرض صفر رد می شود.
	0.1	[1.833512e-13 3.687477e-11]	فرض صفر رد می شود.
	0.2	[2.437635e-13 2.726246e-11]	فرض صفر رد می شود.
	0.3	[3.240388e-13 2.014520e-11]	فرض صفر رد می شود.
	0.4	[4.307084e-13 1.487720e-11]	فرض صفر رد می شود.
	0.5	[5.724567e-13 1.097960e-11]	فرض صفر رد می شود.
	0.6	[7.608336e-13 8.097286e-12]	فرض صفر رد می شود.
	0.7	[1.011205e-12 5.966936e-12]	فرض صفر رد می شود.
	0.8	[1.344020e-12 4.393326e-12]	فرض صفر رد می شود.
	0.9	[1.786506e-12 3.231739e-12]	فرض صفر رد می شود.
	1.0	[2.374926e-12 2.374926e-12]	فرض صفر رد می شود.

شکل ۴: تصمیم‌گیری براساس آزمون نسبت احتمال دنباله ای فازی

$$\tilde{\theta}_0 \approx 0.5, \tilde{\theta}_1 \approx 0.7, \tilde{\alpha} \approx 0.5, \tilde{\beta} \approx 0.1$$

...			
55	0.0	[3.583819e-60 , 1.366995e-50]	فرض صفر رد می شود.
	0.1	[1.067257e-59 , 4.524522e-51]	فرض صفر رد می شود.
	0.2	[3.176356e-59 , 1.494986e-51]	فرض صفر رد می شود.
	0.3	[9.448903e-59 , 4.930157e-52]	فرض صفر رد می شود.
	0.4	[2.809830e-58 , 1.622345e-52]	فرض صفر رد می شود.
	0.5	[8.353708e-58 , 5.325775e-53]	فرض صفر رد می شود.
	0.6	[2.483324e-57 , 1.743720e-53]	فرض صفر رد می شود.
	0.7	[7.382383e-57 , 5.692727e-54]	فرض صفر رد می شود.
	0.8	[2.194941e-56 , 1.852712e-54]	فرض صفر رد می شود.
	0.9	[6.527798e-56 , 6.009408e-55]	فرض صفر رد می شود.
	1.0	[1.942147e-55 , 1.942147e-55]	فرض صفر رد می شود.
...			
100	0.0	[1.187863e-102 , 2.480759e-85]	فرض صفر رد می شود.
	0.1	[8.440136e-102 , 3.322798e-86]	فرض صفر رد می شود.
	0.2	[5.990392e-101 , 4.436858e-87]	فرض صفر رد می شود.
	0.3	[4.247983e-100 , 5.903623e-88]	فرض صفر رد می شود.
	0.4	[3.010447e-99 , 7.824387e-89]	فرض صفر رد می شود.
	0.5	[2.132547e-98 , 1.032490e-89]	فرض صفر رد می شود.
	0.6	[1.510372e-97 , 1.355934e-90]	فرض صفر رد می شود.
	0.7	[1.069759e-96 , 1.771405e-91]	فرض صفر رد می شود.
	0.8	[7.578841e-96 , 2.301079e-92]	فرض صفر رد می شود.
	0.9	[5.371965e-95 , 2.970872e-93]	فرض صفر رد می شود.
	1.0	[3.810445e-94 , 3.810445e-94]	فرض صفر رد می شود.
...			
200	0.0	[3.252771e-213 , 1.965554e-178]	فرض صفر رد می شود.
	0.1	[1.696600e-211 , 3.526336e-180]	فرض صفر رد می شود.
	0.2	[8.829790e-210 , 6.287339e-182]	فرض صفر رد می شود.
	0.3	[4.587381e-208 , 1.113148e-183]	فرض صفر رد می شود.
	0.4	[2.380240e-206 , 1.955313e-185]	فرض صفر رد می شود.
	0.5	[1.234001e-204 , 3.404769e-187]	فرض صفر رد می شود.
	0.6	[6.395079e-203 , 5.872088e-189]	فرض صفر رد می شود.
	0.7	[3.314433e-201 , 1.002193e-190]	فرض صفر رد می شود.
	0.8	[1.718709e-199 , 1.691137e-192]	فرض صفر رد می شود.
	0.9	[8.921185e-198 , 2.818925e-194]	فرض صفر رد می شود.
	1.0	[4.637321e-196 , 4.637321e-196]	فرض صفر رد می شود.

شکل ۵: تصمیم‌گیری براساس آزمون نسبت احتمال دنباله ای فازی

$$\tilde{\theta}_0 \approx 0.5, \tilde{\theta}_1 \approx 0.7, \tilde{\alpha} \approx 0.5, \tilde{\beta} \approx 0.1$$

$$\alpha^U[\circ] = 0/06 \quad \beta^U[\circ] = 0/11$$

$$\alpha^L[\circ] = 0/04 \quad \beta^L[\circ] = 0/09 \quad \text{با در نظر گرفتن}$$

$$\tilde{\alpha}[\circ] = 0/06 \quad \tilde{\beta}[\circ] = 0/11$$

به راحتی می توان مقدار k_0 و k_1 را محاسبه نموده و در مورد رد یا پذیرش هر نمونه اظهار نظر کرد.

داده ها از توزیع نمایی با پارامتر $\theta = 0/6$ ، شبیه سازی شده اند.

اولی در شکل ۳، ۴ و ۵ طراحی شده است، و فرضیه های فازی (۵-۴) را بررسی می کند.

داریم:

(۱) در صورتیکه R_n^L بیشتر یا مساوی $8/545455$ باشد، فرضیه صفر را می پذیریم.

(۲) اگر R_n^L کمتر از $8/545455$ و R_n^U بیشتر از $0/6741573$ باشد، نمونه گیری ادامه پیدا می کند.

(۳) اگر R_n^U کمتر یا مساوی $0/6741573$ باشد، فرضیه صفر رد می شود.

شکل ۳ نشان می دهد، تا زمانی که حداقل ۳ مورد آزمایش نشده باشد، نمی توان در مورد رد یا پذیرش فرضیه صفر اظهار نظر کرد. به ازای $n = 1, 2$ نمی توان تصمیمی

در مورد رد یا پذیرش فرضیه صفر اتخاذ کرد. با توجه به جدول می بینیم که در $n = 3$ و λ برش برابر ۵، می توان در مورد رد فرضیه صفر تصمیم گرفت.

مراجع

- [1] B. Sadeghpour Gildeh, E. Baloui Jamkhaneh, Sequential sampling plan using fuzzy SPRT, Journal of Intelligent Fuzzy Systems 25 (2013) 785–791.
- [2] A. Wald, Sequential analysis, John Wiley, New York, 1947.
- [3] B.F. Arnold, Statistical tests optimally meeting certain fuzzy requirements on the power function and on the sample size, Fuzzy Sets and System 75(2) (1995), 365–372.
- [4] B.F. Arnold, An approach to fuzzy hypotheses testing, Metrika 44 (1996), 119–126.
- [5] B. Sadeghpour Gildeh, Gh. Yari and E. Baloui Jamkhaneh, Acceptance double sampling plan with fuzzy parameter, 11th Joint Conference on Information Sciences, Atlantis Press, China, 2008.
- [6] D.C. Montgomery, Introduction to statistical quality control, Wiley, New York, 1991.
- [7] E. Baloui Jamkhaneh and B. Sadeghpour Gildeh, Sequential sampling plan by variable with fuzzy parameters, Journal of Mathematics and Computer Science 1 (4) (2010), 392–401.
- [8] E. Baloui Jamkhaneh, B. Sadeghpour Gildeh and Gh. Yari, Inspection error and its effects on single sampling plans with fuzzy parameters, Structural and Multidisciplinary Optimization 43 (4) (2011), 555–560.
- [9] E. Baloui Jamkhaneh, B. Sadeghpour Gildeh and Gh. Yari, Acceptance single sampling plan with fuzzy parameter, Iranian Journal of Fuzzy Systems 8 (2) (2011), 47–55.

- [10] H. Torabi and J. Behboodian, Sequential probability ratio test for fuzzy hypotheses testing with vague data, *Austrian Journal of Statistics* 34 (1) (2005), 25–38.
- [11] H. Torabi, J. Behboodian and S.M. Taheri, Neyman-Pearson lemma for fuzzy hypotheses testing with vague data, *Metrika* 64 (3) (2006), 289–304.
- [12] H. Torabi and J. Behboodian, Likelihood ratio test for fuzzy hypotheses testing, *Statistical Papers* 48 (3) (2007), 509–522.
- [13] H. Torabi and S.M. Mirhosseini, Sequential probability ratio tests for fuzzy hypotheses testing, *Applied Mathematical Sciences* 3 (33) (2009), 1609–1618.
- [14] J.J. Buckley, *Fuzzy probability and statistics*, Springer-Verlage Berlin, Heidelberg, 2006.
- [15] K. Dumcic, V. Bahovec and N.K. Zivadinovic, Studying an OC curve of an acceptance sampling: A statistical quality control tool, *Proceeding of the 7th WSEAS International Conference on Mathematics Computers in Business Economics*, cavtat, croatia, JUNE 13-15, 2006, pp. 1–6.
- [16] K. Suey-Sheng, Sequential sampling plan for insect pests, *Phytopathologist Entomologist*, NTU, 1984 , pp. 102–110.
- [17] M.Gh. Akbari, Sequential test of fuzzy hypotheses, *American Open Journal of Statistics* 1 (2011), 87–92 .
- [18] M. Delgado, M.A. Verdegay and M.A. Vila, Testing fuzzy hypotheses: A Bayesian approach, in: M.M. Gupta et al. (Eds), *Approximate Reasoning in Expert Systems*, North-Holland Publishing Co, Amsterdam, 1985, pp. 307–316.
- [19] N.Watanabe and T. Imaizumi, A fuzzy statistical test of fuzzy hypotheses, *Fuzzy Sets and Systems* 53 (1993), 167–178.
- [20] S.M. Taheri and J. Behboodian, A Bayesian approach to fuzzy hypotheses testing, *Fuzzy Sets and Systems* 123 (2001), 39–48.
- [21] S. Hardeo, K. Anwer and A. Muhammad, A visual basic program to determineWald’s sequential sampling plan, *Rev Invest Oper* 25(1) (2004), 84–96.

طرح اصلاحی نمونه‌گیری یک مرحله‌ای متغیر فازی

رباب افشاری^۱، بهرام صادقپور گیلده^۲

^۱دانشگاه فردوسی مشهد، گروه آمار

^۲دانشگاه فردوسی مشهد، گروه آمار

چکیده

یکی از عوامل مهم در پایداری کارخانه‌های تولیدی، میزان کیفیت کالاهای تولیدی است که برای مصرف در اختیار مشتریان قرار می‌گیرد. یکی از روش‌ها در راستای بهبود کیفیت انباشته‌های خروجی، بکارگیری روش بازرسی اصلاحی است. در این مقاله، طرح اصلاحی نمونه‌گیری یک مرحله‌ای فازی برای مشخصه کیفیت متغیر زمانی که از توزیع نرمال تبعیت می‌کند، پیشنهاد می‌شود. نتایج کسب شده بیانگر آن است که کیفیت انباشته خروجی، تحت بازرسی طرح پیشنهادی، همواره بهتر از کیفیت انباشته ورودی است. همچنین به منظور تحلیل و تشریح چگونگی کاربرد طرح پیشنهادی در صنایع تولیدی، مثال کاربردی ارائه می‌شود.

کلمات کلیدی: کنترل کیفیت آماری، طرح نمونه‌گیری یک مرحله‌ای برای پذیرش، متوسط کیفیت خروجی و حساب اعداد فازی.

۱ مقدمه

طرح‌های نمونه‌گیری برای پذیرش، یکی از ابزارهای مهم در کنترل کیفیت آماری جهت اخذ تصمیم درباره انباشته‌ای از محصولات تولیدی است. با توجه به اینکه بتوان مشخصه کیفیت مورد نظر را به صورت یک مقیاس نسبی بیان کرد یا نه، طرح‌های نمونه‌گیری برای پذیرش به ترتیب، طرح نمونه‌گیری متغیر و وصفی نامیده می‌شود. هر طرح نمونه‌گیری برای پذیرش دارای پارامترهایی است و یافتن مقادیر عددی این پارامترها به منزله طراحی آن طرح محسوب می‌شود. طرح نمونه‌گیری یک مرحله‌ای جهت پذیرش انباشته‌ای از تولیدات، به دلیل سادگی در اجرا یکی از پر طرفدارترین روش از سوی صنعتگران و تولیدکنندگان است. چاکرابورتی (۱۹۹۲) یک کلاس از طرح‌های نمونه‌گیری ساده مبتنی بر بهینه‌سازی در محیط فازی را مورد مطالعه قرار داد. طرح‌های نمونه‌گیری جهت پذیرش برای مشخصه‌های کیفیت متغیر در محیط فازی توسط ژورزوسکی (۲۰۰۲) بررسی شد.

^۱رباب افشاری : robab.afshari@mail.um.ac.ir

یکی از شیوه‌های مرسوم جهت ارتقاء کیفیت انباشته‌های خروجی استفاده از روش بازرسی اصلاحی است. داج و رومینگ (۱۹۵۹)، طرح نمونه‌گیری یک مرحله‌ای اصلاحی کلاسیک را معرفی کردند. بالوئی جامخانه و صادقی‌پورگیله (۲۰۱۲)، طرح نمونه‌گیری دو مرحله‌ای اصلاحی وصفی فازی را مطالعه کردند. روش بازرسی اصلاحی تحت اجرای طرح نمونه‌گیری تعویقی چندگانه بیزی برای مشخصه کیفیت وصفی با توزیع پیشین گاما توسط لاتا و سایا (۲۰۱۶) بررسی شد. افشاری و صادقی‌پورگیله (۲۰۱۸)، طرح نمونه‌گیری تعویقی چندگانه فازی متغیر و طرح نمونه‌گیری یک مرحله‌ای فازی متغیر^۱ (FSSP) را معرفی و عملکرد آنها را در مقایسه با هم بررسی کردند. در مقاله حاضر، طرح FSSP در حضور بازرسی اصلاحی، زمانی که مشخصه کیفیت مورد نظر دارای توزیع نرمال با حد مشخصه بالایی بوده و نسبت اقلام معیوب^۲ (p) انباشته کمیتی فازی باشد، پیشنهاد می‌شود. در بخش بعد الگوریتم اجرای طرح پیشنهادی معرفی و معیار متوسط کیفیت خروجی در بخش سوم بررسی می‌شود. چگونگی کاربرد روش پیشنهادی بر روی داده‌های واقعی، در بخش چهارم آورده می‌شود. بخش پایانی به بیان نتایج اخذ شده اختصاص دارد.

۲ الگوریتم اجرای طرح اصلاحی نمونه‌گیری یک مرحله‌ای فازی متغیر

در برنامه‌های نمونه‌گیری برای پذیرش، معمولاً در صورت رد انباشته نیاز است تا اقدام اصلاحی انجام پذیرد. این برنامه‌های نمونه‌گیری را به علت تأثیری که فعالیت‌های بازرسی بر کیفیت نهایی محصول خروجی دارد برنامه‌های بازرسی اصلاحی می‌نامند. در این بخش طرح اصلاحی FSSP متغیر را در حالتی که مقدار p نادقیق است پیشنهاد می‌کنیم. به منظور سادگی در محاسبات، فرض کنید نسبت اقلام معیوب انباشته عدد فازی مثلثی $\tilde{p} = (a_1, a_2, a_3)$ بوده و اندازه انباشته برابر N باشد. همچنین فرض کنید مشخصه کیفیت مورد نظر دارای حد مشخصه بالایی U بوده و از توزیع نرمال با میانگین نامعلوم μ و انحراف استاندارد σ پیروی کند. اگر انحراف استاندارد جامعه σ معلوم باشد، روش اجرای طرح اصلاحی FSSP با انحراف استاندارد معلوم^۳ (KSD-FSSP) برای تصمیم‌گیری درباره انباشته به صورت زیر پیشنهاد می‌شود:

گام اول نمونه‌ای به اندازه n_σ مانند $(X_1, X_2, \dots, X_{n_\sigma})$ از انباشته جمع‌آوری و مقدار آماره $V_\sigma = \frac{U - \bar{X}}{\sigma}$ را به دست می‌آوریم، که در آن $\bar{X} = \frac{1}{n_\sigma} \sum_{i=1}^{n_\sigma} X_i$ است.

گام دوم اگر $v_\sigma \geq k_\sigma$ آنگاه انباشته را پذیرفته و اقلام معیوب مشاهده شده در نمونه را با اقلام سالم جایگزین می‌کنیم.

اگر $v_\sigma < k_\sigma$ آنگاه انباشته را رد کرده و غربالگری صد در صد روی انباشته انجام می‌دهیم. یعنی اینکه، انباشته مورد بازرسی صد در صد قرار گرفته و تمامی اقلام معیوب مشاهده شده با اقلام سالم جایگزین می‌شوند. v_σ مقدار مشاهده شده آماره V_σ و k_σ عدد پذیرش تحت اجرای طرح اصلاحی KSD-FSSP است.

بنابراین طرح اصلاحی KSD-FSSP دارای دو پارامتر k_σ و n_σ است.

با استفاده از روش باکلی (۲۰۰۶)، λ -برش احتمال فازی پذیرش در طرح KSD-FSSP برابر است با:

$$\begin{aligned} \tilde{P}_{a\sigma}(\tilde{p})[\lambda] &= \left\{ \Phi((z_p - k_\sigma)\sqrt{n_\sigma}) \mid p \in \tilde{p}[\lambda] \right\}, \\ &= [P_{a\sigma}^-(\lambda), P_{a\sigma}^+(\lambda)], \end{aligned} \quad (1)$$

که در آن $P_{a\sigma}^-(\lambda)$ و $P_{a\sigma}^+(\lambda)$ به ترتیب می‌نیم و ماکزیم $\tilde{P}_{a\sigma}(\tilde{p})[\lambda]$ هستند. z_p چندک مرتبه $(1-p)$ ام توزیع نرمال استاندارد بوده و برابر $z_p = \Phi^{-1}(1-p)$ است.

به طور مشابه، در حالتی که انحراف استاندارد جامعه نامعلوم باشد، می‌توان روش اصلاحی پیشنهاد شده را در طرح FSSP با انحراف استاندارد نامعلوم^۴ (USD-FSSP) به کار گرفت. با این تفاوت که وقتی انحراف استاندارد جامعه نامعلوم است، در گام اول، از انحراف استاندارد نمونه‌ای $S = \sqrt{\frac{1}{n_s-1} \sum_{i=1}^{n_s} (X_i - \bar{X})^2}$

^۱Fuzzy single sampling plan

^۲Production of nonconforming items

^۳Known standard deviation FSSP

^۴Unknown standard deviation FSSP

به عنوان برآورد σ استفاده می‌شود که در آن n_s ، اندازه نمونه گرفته شده از انباشته است. به علاوه، در حالت انحراف استاندارد نامعلوم از نمادهای k_s, v_s, V_s به ترتیب به جای $k_\sigma, v_\sigma, V_\sigma$ استفاده می‌کنیم. اگر چه الگوریتم اجرای طرح پیشنهادی FSSP در هر دو حالت انحراف استاندارد معلوم و نامعلوم، مشابه یکدیگر هستند اما محاسبه احتمالات پذیرش در تعیین پارامترهای طرح USD-FSSP کمی متفاوت با طرح KSD-FSSP است. مطابق دانکن (۱۹۸۶)، داریم:

$$\bar{X} + k_s S \sim N\left(\mu + k_s \sigma, \frac{\sigma^2}{n_s} + k_s^2 \frac{\sigma^2}{v n_s}\right). \quad (2)$$

با توجه به رابطه (۲) و با بکارگیری روش باکلی، λ -برش احتمال فازی پذیرش در طرح USD-FSSP به صورت زیر محاسبه می‌شود:

$$\begin{aligned} \tilde{P}_{as}(\tilde{p})[\lambda] &= \left\{ \Phi\left((z_p - k_s) \sqrt{\frac{n_s}{1 + k_s^2/v}}\right) \mid p \in \tilde{p}[\lambda] \right\}, \\ &= [P_{as}^-(\lambda), P_{as}^+(\lambda)]. \end{aligned} \quad (3)$$

۳ متوسط کیفیت خروجی فازی در طرح اصلاحی FSSP

اگر انباشته‌ای از محصولات یک فرآیند تولیدی، که دارای نسبت اقلام معیوب p است، تهیه و بازرسی اصلاحی در آن انجام پذیرد، آنگاه سطح متوسط کیفیت انباشته که در بلندمدت حاصل می‌شود را متوسط کیفیت خروجی ^۵ (AOQ) می‌نامند. در این بخش، هدف ما یافتن رابطه‌ای برای محاسبه AOQ در طرح اصلاحی پیشنهادی FSSP در هر دو حالت انحراف استاندارد معلوم و نامعلوم است. در ادامه خواهیم دید که اگر نسبت اقلام معیوب انباشته نادقیق گزارش شده باشد، آنگاه مقدار AOQ نیز نادقیق خواهد بود. بنابراین این مقدار را متوسط کیفیت خروجی فازی نامیده و با نماد $\widetilde{AOQ}$ نشان می‌دهیم.

قضیه ۱۰۳. طرح اصلاحی KSD-FSSP با پارامترهای n_σ و k_σ را در نظر بگیرید. فرض کنید اندازه انباشته و نسبت اقلام معیوب به ترتیب برابر N و $\tilde{p}$ باشند. اگر مشخصه کیفیت مورد نظر دارای توزیع نرمال $N(\mu, \sigma^2)$ (معلوم σ) با حد مشخصه بالایی U باشد، آنگاه به ازای هر $\lambda \in [0, 1]$ ، λ -برش مقدار $\widetilde{AOQ}_\sigma$ برابر است با:

$$\begin{aligned} \widetilde{AOQ}_\sigma[\lambda] &= \left\{ \frac{(N - n_\sigma)p P_{a\sigma}}{N}, p \in \tilde{p}[\lambda] \right\} \\ &= [AOQ_\sigma^-(\lambda), AOQ_\sigma^+(\lambda)], \end{aligned} \quad (4)$$

که در آن $P_{a\sigma} = \Phi((z_p - k_\sigma) \sqrt{n_\sigma})$ است.

اثبات. نمونه‌ای به حجم n_σ از انباشته تحت بازرسی انتخاب می‌کنیم. پس از انجام بازرسی دو حالت ممکن است در خصوص تصمیم‌گیری درباره انباشته رخ دهد. یا انباشته بعد از بازرسی با احتمال فازی $\tilde{P}_{a\sigma}$ پذیرفته می‌شود و یا با احتمال $1 - \tilde{P}_{a\sigma}$ رد می‌شود. اگر انباشته رد شود آنگاه انباشته تحت بازرسی 100% قرار گرفته و همه اقلام معیوب مشاهده شده با اقلام سالم جایگزین می‌شوند. در نتیجه تعداد اقلام معیوب موجود در انباشته خروجی برابر صفر خواهد شد. اگر انباشته پذیرفته شود در آن صورت فقط اقلام معیوب مشاهده شده در نمونه به اندازه n_σ با اقلام سالم جایگزین می‌شوند. بنابراین تعداد اقلام معیوب مورد انتظار در واحدهای باقی‌مانده که بدون بازرسی در اختیار مشتری قرار می‌گیرد برابر با $(N - n_\sigma)\tilde{p}$ است. در نتیجه در این فرآیند بازرسی، با استفاده از تعریف میانگین فازی و بکارگیری روش باکلی $\square$ اثبات قضیه کامل می‌شود. (۲۰۰۶)

فرض کنید انحراف استاندارد جامعه نامعلوم باشد. مدل ریاضی برای محاسبه AOQ تحت اجرای طرح اصلاحی USD-FSSP که با نماد $\widetilde{AOQ}_s$ نشان می‌دهیم، دقیقاً مشابه با حالتی است که در حالت انحراف استاندارد معلوم در قضیه ۱۰۳ بیان شد، تنها با این تفاوت که در این حالت، از انحراف استاندارد نمونه‌ای به عنوان برآورد σ استفاده می‌شود.

^۵ Average outgoing quality

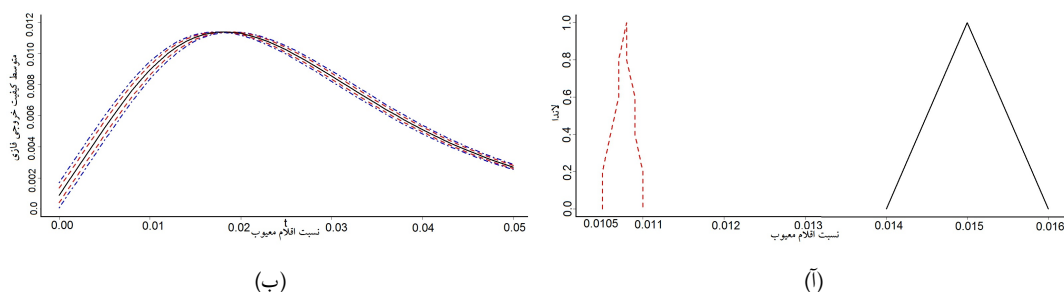
جدول ۱: داده‌های میزان نمک موجود در کشک تولیدی شرکت پگاه استان زنجان بر حسب درصد

۱/۲۳	۱/۲۵	۱/۲۰	۱/۱۰	۱/۳۵	۱/۴۰	۱/۲۳	۱/۲۵	۱/۲۰	۱/۱۰	۱/۳۵	۱/۴۰	۱/۲۳	۱/۲۵
۱/۲۰	۱/۱۰	۱/۳۵	۱/۴۰	۱/۲۵	۱/۲۵	۱/۳۵	۱/۳۶	۱/۵۰	۱/۳۵۰	۱/۳۶	۱/۳۶	۱/۵۰	۱/۳۵۰
۱/۲۸	۱/۲۸	۱/۲۸	۱/۲۸	۱/۲۷	۱/۲۷	۱/۴۰	۱/۴۰	۱/۲۰	۱/۲۰	۱/۲۰	۱/۴۶	۱/۳۰	۱/۴۶
۱/۳۰	۱/۴۶	۱/۳۰	۱/۳۵	۱/۴۵	۱/۳۰	۱/۳۰	۱/۲۰	۱/۳۵	۱/۴۵	۱/۳۰	۱/۳۰	۱/۲۰	۱/۳۵
۱/۴۵	۱/۳۰	۱/۳۰	۱/۲۰	۱/۴۵	۱/۲۳	۱/۲۳							

۴ مثال کاربردی

در این بخش، طرح اصلاحی نمونه‌گیری یک مرحله‌ای پیشنهاد شده در محیط فازی، بر روی داده‌های واقعی به کار گرفته می‌شود. به همین منظور، مجموعه داده‌های میزان نمک موجود در کشک تولید شده توسط شرکت پگاه استان زنجان در نظر گرفته شده است. هدف، تصمیم‌گیری درباره رد یا پذیرش انباشته‌ای از تولیدات کشک بر اساس طرح اصلاحی پیشنهاد شده FSSP است. در این تصمیم‌گیری، مشخصه کیفیت مورد نظر، میزان نمک موجود در محصول تولیدی بوده و از نوع متغیر می‌باشد. حداکثر میزان مجاز نمک در این محصول از سوی استاندارد ملی ایران، مقدار دو درصد در نظر گرفته شده است ($U = 0.02$). فرض کنید طبق معاهده بسته شده بین مصرف‌کننده و تولیدکننده، مقادیر $\overline{AQL}$ و $\overline{LQL}$ به ترتیب تقریباً برابر 0.01 و 0.045 با احتمالات پذیرش به ترتیب تقریباً 0.95 و 0.1 باشند. در این صورت با توجه به مقادیر گزارش شده در جدول ۳،۳ در افشاری (۲۰۱۸)، پارامترهای طرح برابر $n_s = 63$ و $k_s = 1/97$ است. چنانچه اشاره شد، یکی از فرض‌های اساسی در بکارگیری طرح پیشنهادی USD-FSSP برای پذیرش انباشته‌ای از تولیدات، نرمال بودن مشخصه کیفیت مورد نظر است. بنابراین ابتدا باید تعیین کنیم که آیا داده‌های ثبت شده دارای توزیع نرمال هستند یا نه. آزمون شاپیرو-ویلک نشان داد که $p\text{-value} = 0.085$ است که مقداری بزرگتر از 0.05 است. یعنی این که نرمال بودن توزیع داده‌ها رد نمی‌شود. اما به دلیل محدودیت در تکرار آزمایشات و نمونه‌برداری، فرضیه نرمال بودن را می‌پذیریم. بنابراین جهت تصمیم‌گیری درباره این انباشته مبتنی بر طرح پیشنهاد شده USD-FSSP نمونه‌ای به اندازه ۶۳ از انباشته جمع‌آوری می‌کنیم. این داده‌ها در جدول ۱ آورده شده است. فرض کنید اندازه انباشته برابر $N = 500$ باشد. مقادیر میانگین نمونه‌ای و انحراف استاندارد نمونه‌ای به ترتیب برابر $\bar{x} = 1/305$ و $s = 0.096$ به دست می‌آید. در نتیجه داریم $v_s = \frac{2-1/305}{0.096} = 7/23 \geq 1/97$ چون $v_s \geq 1/97$ پس انباشته تولیدی پذیرفته می‌شود و تنها اقلام ناسالم مشاهده شده در نمونه با سالم جایگذاری و تحویل مشتری می‌گردد. به عنوان مثال فرض کنید نسبت اقلام معیوب انباشته برابر $(0.014, 0.015, 0.016)$ باشد. در این صورت مطابق قضیه ۱۰۳ در صورتی که انحراف استاندارد جامعه نامعلوم باشد، داریم $AOQ_s[0] = [0.0105, 0.0110]$ و $AOQ_s[1] = 0.0108$ یعنی اینکه هرگاه احتمال معیوب بودن قطعه‌ای در یک انباشته برابر تقریباً 0.015 باشد، آنگاه متوسط نسبت اقلام معیوب انباشته پس از بازرسی اصلاحی تحت اجرای طرح پیشنهادی، برابر تقریباً 0.0108 بوده و به وضوح کاهش یافته است. مطابق شکل ۱(آ) مشاهده می‌شود $AOQ \leq \tilde{p}$ که در آن نماد $\leq$ به مفهوم کمتر یا مساوی است. برای ترسیم منحنی AOQ در طرح اصلاحی پیشنهادی، ساختار $\tilde{p}$ را به صورت $\tilde{p}_t = (t, b_2 + t, b_3 + t)$ بازنویسی می‌کنیم که در آن $b_i = a_i - a_1$ ($i = 2, 3$) و $t \in [0, 1 - b_3]$ است. فرض کنید $\tilde{p} = (0.0001, 0.0002, 0.0003)$ باشد آنگاه $\tilde{p}_t = (t, t + 0.0001, t + 0.0002)$ باشد. مطابق شکل ۱(ب)، مشاهده می‌شود که منحنی AOQ به صورت نواری ظاهر شده است که دارای کران‌های بالا و پایین است به طوری که این کران‌ها، با کاهش میزان ابهام در مقدار نسبت اقلام معیوب انباشته، به همدیگر نزدیک شده و حالت کلاسیک رخ می‌دهد. بعلاوه به آسانی می‌توان دید اگر نسبت اقلام معیوب ورودی خوب باشد، آنگاه متوسط نسبت اقلام معیوب خروجی نیز خوب است. همچنین، مشاهده می‌شود هرگاه نسبت اقلام معیوب ورودی بد باشد آنگاه انباشته خروجی از متوسط نسبت اقلام معیوب خوب برخوردار است. زیرا در این حالت، انباشته تحت بازرسی، رد شده و بازرسی ۱۰۰٪ انجام می‌گیرد. در نتیجه تمام اقلام معیوب مشاهده شده با اقلام سالم جایگزین می‌شوند. در بین این دو نقطه حدی، AOQ به حداکثر مقدار خود می‌رسد و سپس کاهش می‌یابد. حداکثر مقدار AOQ نشان‌دهنده بیش‌ترین متوسط نسبت اقلام معیوب در انباشته‌های خروجی پس از بازرسی اصلاحی است که اصطلاحاً حد متوسط کیفیت خروجی فازی نامیده و با نماد $\overline{AOQL}$ نشان داده می‌شود. با توجه به شکل ۱(ب)، $\overline{AOQL}$ تقریباً برابر 0.0114 است. به عبارت دیگر به هر میزان که نسبت اقلام معیوب انباشته‌های ورودی بد باشد هیچگاه میزان متوسط نسبت اقلام معیوب در انباشته‌های خروجی تحت بازرسی اصلاحی در طرح USD-FSSP از

تقریباً ۱/۱۴٪ افزایش خواهد یافت.



شکل ۱: (ا) نمودار توابع عضویت $\bar{p}$ و $\widetilde{AOQ}$ (ب) منحنی $\widetilde{AOQ}$ به ازای $\lambda = 0, 0.4, 1$

بحث و نتیجه‌گیری

در مطالعه حاضر، به منظور ارتقاء و بهبود کیفیت انباشته‌ای از تولیدات، طرح اصلاحی FSSP برای مشخصه کیفیت متغیر نرمال زمانی که پارامتر نسبت کمیتی فازی باشد، معرفی شد. سپس با بکارگیری مفهوم میانگین فازی، رابطه ریاضی جهت محاسبه معیار متوسط کیفیت خروجی فازی به دست آمد. نتایج کسب شده نشان دادند که این معیار، به دلیل فازی بودن کیفیت فرآیند، خود نیز کمیتی فازی به دست می‌آید. به منظور تفهیم بهتر روش اجرای طرح معرفی شده، از برخی مثال‌های عددی و کاربردی بهره گرفته شد. نتایج به دست آمده حاکی از آن است که بازرسی اصلاحی در حضور پارامتر نسبت فازی، باعث بهبود کیفیت انباشته‌های خروجی می‌گردد و بنابراین بکارگیری آن می‌تواند در جذب مصرف‌کنندگان بیشتر موثر واقع شود. به علاوه نشان داده شد روش معرفی شده تعمیمی از روش مذکور در حالت کلاسیک است.

مراجع

[۱] افشاری، ر. (۱۳۹۶). مباحثی در طرح‌های نمونه‌گیری شرطی برای پذیرش، رساله دکتری، گروه آمار، دانشگاه فردوسی مشهد.

- [2] Afshari, R. and Sadeghpour Gildeh, B. (2018). Fuzzy multiple deferred state variable sampling plan, *J. of Intelligent and Fuzzy Sys.*, In Press, DOI: 10.3233/JIFS-17907.
- [8] Buckley, J. J., (2006). *Fuzzy probability and statistics*, Springer-Verlag, Berlin Heidelberg.
- [4] Baloui Jamkhaneh, E. and Sadeghpour Gildeh, B. (2012). Acceptance double sampling plan using fuzzy poisson distribution, *World Applied Sciences J.*, 16(11), 1578-1588.
- [5] Chakraborty, T. K. (1992). A class of single sampling plans based on fuzzy optimization. *Quality Control and Applied Statistics*, 37(7), 359-362.
- [6] Dodge, H. F. and Roming, H. G. (1959). *Sampling inspection tables*, 2th ed, Wiley, NewYork.
- [7] Duncan, A. J. (1986). *Quality Control and Industrial Statistics*, 5th ed. Homewood, IL: Irwin.

- [8] Grzegorzewski, P. (2002). Acceptance sampling plans by variables for vague data, *Soft Methods in Probability, Statistics and Data Analysis*. Physica-Verlag HD.
- [9] Latha, M. and Subbiah, K. (2016). Selection of bayesian multiple deferred state (bmbs-1) sampling plan based on quality regions. *Int. J. of Recent Scientific Research*, 6(5), 3864-3867.

کاربرد رگرسیون لجستیک فازی در مدل سازی شدت اختلال اوتیسم

علی بهنام پور^۱، اکبر بیگلریان^۲، عنایت الله بخشی^۳، مجید قربانی^۴، مهشید نامداری^۵

^۱گروه آمارزیستی، دانشگاه علوم بهزیستی و توانبخشی

^۲گروه آمارزیستی، دانشگاه علوم بهزیستی و توانبخشی

^۳گروه آمارزیستی، دانشگاه علوم بهزیستی و توانبخشی

^۴پیش دبستانی و دبستان غیردولتی اختلال رفتاری-هیجانی و طیف اوتیسم پرورش

^۵گروه آمارزیستی، دانشگاه علوم پزشکی شهید بهشتی

چکیده

در این مطالعه از مدل رگرسیون لجستیک فازی برای دستیابی به مدلی به منظور تعیین شدت اختلال اوتیسم استفاده شده و نتایج آن مورد بررسی قرار گرفته است. برای برآزش مدل از خرده مقیاس های حرکات کلیشه ای، ارتباط و تعامل اجتماعی به عنوان متغیر های ورودی و همچنین نظر درمانگر در خصوص شدت اختلال به عنوان متغیر خروجی فازی، استفاده شده است. برآورد ضرایب به دو روش کمترین مربعات و کمترین قدرمطلق اختلافات انجام شده و معیار نیکویی برآزش نیز محاسبه شده است. لگاریتم بخت امکانی برای یک فرد به صورت فرضی نیز محاسبه شد.

کلمات کلیدی: رگرسیون لجستیک فازی، بخت امکانی، مولفه زبانی، اوتیسم، اختلال در خودماندگی

۱ مقدمه

اختلال طیف اوتیسم^۱ (ASD) یکی از اختلالات عصب رشدی^۲ است که با آسیب در تعاملات اجتماعی و رفتارها، علایق و حرکات کلیشه ای و تکراری شناخته می شود (۳). این اختلال به عنوان وخیم ترین و در عین حال، ناشناخته ترین اختلال دوران کودکی شناخته شده است (۱۶). برای اولین بار در سال ۱۹۴۳، یک روانپزشک اتریشی، به نام لئو کانر^۳، در مقاله معروف خود، به توصیف و معرفی این کودکان پرداخت (۸). آنچه او مطالعه کرده بود کودکانی بودند که از سال اول زندگی شان به

^۱علی بهنام پور: al.behnampour@uswr.ac.ir

^۱Autism Spectrum Disorder

^۲Neurodevelopmental Disorder

^۳Leo Kanner

طور مداوم از هرگونه تماس و ارتباطی اجتناب می کردند و ویژگی هایی همچون درخودماندگی، تکرار یا پژواک کلامی، اصرار بر یک نواختی و فقدان تماس چشمی داشتند (۱۷).

شیوع این اختلال روز به روز در حال افزایش است. مطابق آمار مرکز کنترل و پیشگیری بیماری های آمریکا^۴ (CDC) شیوع این اختلال در سال ۲۰۰۲، یک در ۱۵۰ نفر، در سال های ۲۰۰۴ تا ۲۰۰۶، یک در ۱۱۰ (۱۳) در سال ۲۰۰۸، یک در ۸۸ نفر (۴) و در سال ۲۰۱۴، یک در ۶۸ نفر (۵) گزارش شده است. مطالعات شیوع شناسی در ایران نیز آمار رو به افزایشی از ابتلا به اوتیسم در کودکان را نشان می دهند. برای مثال، در مطالعه (۶) میزان شیوع این اختلال را در اصفهان ۱۵/۱۲ و در شهرکرد ۹۷/۹ در هر ۱۰ هزار کودک گزارش کرده اند. صمدی و همکاران نیز شیوع اوتیسم را در ایران در سال ۲۰۰۷، ۲۶/۶ در ۱۰ هزار کودک (۱۴) و در سال ۲۰۱۴، ۲/۹۵ در ۱۰ هزار کودک گزارش کرده اند (۱۵). با توجه به مشکلات زیادی که این اختلال برای کودک، خانواده و جامعه به وجود می آورد و نیز افزایش روز افزون آن، ضرورت غربالگری^۵ و تشخیص^۶ زود هنگام و ارائه مداخلات بهنگام اهمیت ویژه ای دارد (۲).

از سال ۱۹۸۰ که برای اولین بار اختلال اوتیسم به صورت یک طبقه جداگانه شناخته شد، تا سال ۲۰۱۳ که پنجمین ویرایش راهنمای تشخیصی و آماری اختلالات روانی^۷ (DSM-5) (۳) از سوی انجمن روانپزشکی آمریکا منتشر شد، تغییر و تحولات وسیعی در حوزه اوتیسم صورت گرفته است. یکی از انتقادات عمده ای که به نسخه های قبلی DSM وارد می شد، این بود که در آنها طبقه بندی تفکیکی (افتراقی) شفافیت نداشت. برای مثال، تحقیقات نشان می دهند که در مورد بسیاری از اختلالات، تمایز بین نشانه های عمده و نشانه های خفیف تر، به «درجه» ربط دارد، نه به «نوع» (۷). به همین خاطر و به دلیل مشکلاتی که در تفکیک اختلال اوتیسم، اختلال آسپرگر^۸، اختلال رت^۹ و اختلال فروپاشی کودکی^{۱۰} وجود داشت، DSM-5 همه این اختلالات را حذف و آنها را تحت لوای «اختلال طیف درخودماندگی (اوتیسم)» معرفی کرده است (۳). پنجمین ویرایش تشخیصی و آماری اختلالات روانی سه سطح شدت برای اوتیسم معرفی کرده است: سطح ۱: نیازمند به حمایت؛ سطح ۲: نیازمند به حمایت زیاد؛ و سطح ۳: نیازمند به حمایت بسیار زیاد.

پروفسور زاده در مقاله سال ۱۹۶۹ خود گفته است: «پیشگیری سیستم های بیولوژیکی ما را مجبور می کند تا برای تجزیه و تحلیل آن ها جایگزینی برای روش های تحلیلی گذشته خود بیابیم. باید بپذیریم که درجه ای از ابهام در رفتار سیستم های بیولوژیکی و ماهیت آن ها وجود دارد» (۱۹). کمی بعد دوباره بر این نکته تاکید داشته و می گویند: «اصلی ترین کاربرد متغیر های زبانی و الگوریتم های فازی در علوم از قبیل اقتصاد، مدیریت، هوش مصنوعی، روانشناسی، زبان شناسی، پزشکی و بیولوژی است» (۲۰).

برای مدل سازی روابط بین متغیرهایی که مشاهدات آن ها نادقیق است نمی توان مدل های معمول آماری را که بر اساس مشاهدات دقیق و برخی فرضیات توزیعی استوار هستند به کار برد. برای مدلسازی، توصیف و تحلیل روابط بین این گونه متغیرها می توان از مجموعه های فازی استفاده کرد. از زمان ارائه ی نظریه مجموعه های فازی، استفاده از این نظریه در بخش های وسیعی از علم آمار که کاربردهای فراوانی از آن ها در تحقیقات کاربردی پزشکی متصور است در حال گسترش است، که از جمله می توان به خوشه بندی فازی، تحلیل ممیزی فازی، رگرسیون فازی، تشخیص بیماری ها با رویکرد فازی و مدلسازی آشکار شدن و انتقال بیماری هایی با دوره پنهان اشاره کرد.

تقریباً اکثر مطالعات انجام شده در زمینه رگرسیون فازی، در زمینه مدل های خطی، صورت گرفته است و مدل های غیر خطی به ندرت در محیط فازی بررسی شده اند. یکی از مشهور ترین مدل های غیر خطی، مدل رگرسیون لوجستیک است. این مدل در علوم پزشکی و مطالعات اجتماعی کاربرد فراوان دارند. رگرسیون لوجستیک روش مفیدی است برای توصیف رابطه ی بین یک یا چند فاکتور و متغیری که مقادیری به صورت کیفی و گسسته اختیار می کند. اما در عمل شرایطی پیش می آید که مفروضات رگرسیون لوجستیک آماری نقض شده یا قابل بررسی نیست و لذا استفاده از آن با اشکال مواجه می شود. به طور مثال ممکن است مقادیر متغیر پاسخ به جای عدد به صورت مولفه های زبانی بیان شود. اینگونه مشاهدات نادقیق در تحقیقات بالینی زیاد به چشم می خورند. در حالی که در مدل رگرسیون لوجستیک آماری برای برازش مدل از مشاهدات دقیق استفاده می شود، در چنین مواردی می توان از رگرسیون لوجستیک فازی استفاده کرد (۹).

^۴Centers for Disease Control and Prevention

^۵Screening

^۶Diagnosis

^۷Diagnostic and Statistical Manual of Mental Disorders-Fifth Edition

^۸Asperger disorder

^۹Rett syndrome

^{۱۰}Childhood disintegrative disorder

تشخیص اختلال طیف اوتیسم بعضاً می‌تواند بسیار دشوار باشد زیرا هیچ تست پزشکی دقیق مانند تست خون برای تشخیص این نوع اختلالات وجود ندارد. به منظور تشخیص اختلالات پزشکان و درمانگران به رفتار و رشد کودک توجه کرده و در آنها به دنبال نشانه‌هایی مرتبط با اختلال می‌گردند. این اختلال می‌تواند گاهی در ۱۸ ماهگی یا قبل از آن تشخیص داده شود. در ۲ سالگی تشخیص توسط متخصص حرفه‌ای می‌تواند بسیار قابل اطمینان تلقی شود. با این حال بسیاری از کودکان تا سن بالاتر تشخیص نهایی را دریافت نمی‌کنند. این تأخیر به این معناست که کودکان مبتلا به اوتیسم ممکن است کمکی که لازم دارند را دریافت نکنند. اختلاف نظر درمانگران در تشخیص وجود و یا حتی تعیین شدت این اختلال، همچنین استفاده از مولفه‌های زبانی به نوعی عدم قطعیت در نتیجه گزارش شده منجر خواهد شد. لذا استفاده از رگرسیون لوجستیک فازی به جای رگرسیون لوجستیک توصیه می‌شود.

۲ روش کار

این مطالعه از نوع تحلیلی می‌باشد که جمع‌آوری اطلاعات بیماران بصورت مقطعی صورت گرفته است. جامعه آماری افراد مبتلا به اختلال طیف درخودماندگی هستند که به پیش دبستانی و دبستان غیردولتی اختلال رفتاری-هیجانی و طیف اوتیسم پرورش شهرگران مراجعه نموده‌اند. نمونه پژوهش ۲۰ کودک مبتلا به این اختلال بودند که به صورت تمام شماری انتخاب شدند. از والدین و یا درمانگر آنها خواسته شد تا آزمون تشخیصی اوتیسم گیلیام - نسخه فارسی را به دقت برای این کودکان تکمیل نمایند. نمره ی خام مربوط به هر یک از خرده مقیاس‌های حرکات کلیشه‌ای، ارتباط و تعامل اجتماعی به عنوان متغیرهای ورودی دقیق مورد استفاده قرار گرفتند. همچنین نظر درمانگر دیگری که افراد را مورد ارزیابی قرار داده بود، در مورد شدت اختلال فرد به عنوان متغیر پاسخ فازی در نظر گرفته شد. از درمانگر خواسته شد تا نظر خود را در مورد شدت اختلال افراد بصورت واژه‌های بیانی (درجه خفیف، متوسط و شدید) ابراز نماید. در نمونه جمع‌آوری شده ۲ کودک به دلیل نداشتن توانایی تکلم، نمره خرده مقیاس ارتباط برای آنها، و برای یک نفر هم نمره خرده مقیاس تعامل اجتماعی ثبت نگردید. لذا تعداد افراد مورد مطالعه به ۱۷ نفر کاهش یافت. توجه محققان نسبت به درمان موجب شده است که آنها به دنبال ابزارهای دقیقی برای تشخیص اختلال اوتیسم باشند. در این راه ابزارهای تشخیصی متعددی تهیه شده است که از مهم‌ترین آنها می‌توان به آزمون گارز (GARS) اشاره کرد. این آزمون به عنوان یک آزمون معتبر توسط گیلیام در سال ۱۹۹۴ تهیه شده است. پایایی گارز در دامنه قابل پذیرش، پذیرفته شده است. مطالعات انجام شده نمایانگر ضریب آلفای ۰/۹۰ برای رفتارهای کلیشه‌ای، ۰/۸۹ برای ارتباط، ۰/۹۳ برای تعامل اجتماعی، ۰/۸۸ برای اختلالات رشدی و ۰/۹۶ در نشانه‌شناسی اوتیسم است (۱). ابزار مورد استفاده در این مطالعه آزمون تشخیصی اوتیسم گیلیام - نسخه فارسی می‌باشد.

در این مطالعه با استفاده از مدل رگرسیون لوجستیک فازی که در (۱۱) پیشنهاد شد، مدل فازی بر اساس امکان موفقیت، که امکان‌ها به صورت واژه‌های بیانی مشخص شده‌اند، معرفی شده و سپس "بخت امکانی" بر اساس متغیرهای مستقل مشاهده شده که مقادیر دقیق داشتند، مدل سازی شد. در این مطالعه به دنبال برازش مدلی برای تعیین شدت اختلال اوتیسم با استفاده از رگرسیون لوجستیک فازی هستیم. برای ارزیابی مدل شاخصی به نام "شاخص قابلیت بر اساس فاصله‌های فازی"^{۱۱} محاسبه شده است.

رگرسیون لوجستیک فازی:

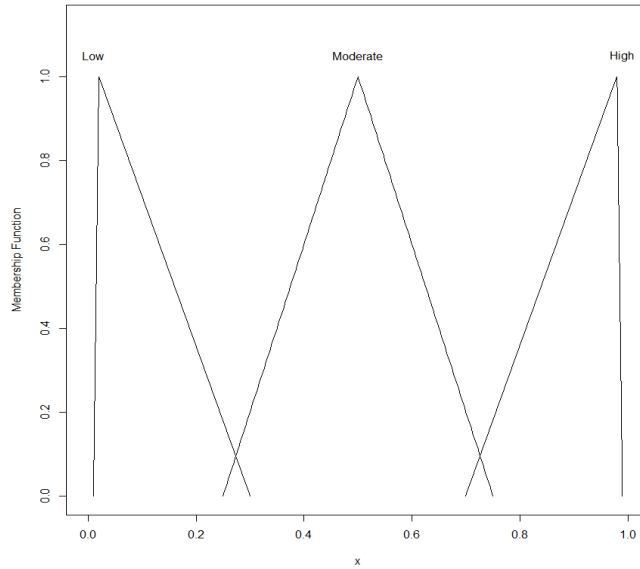
اگر درجه سازگاری با شرایط معلومی را برای فرد i ام با μ_i ($0 < \mu_i < 1$) نشان دهیم، نسبت $\frac{\mu_i}{1-\mu_i}$ را بخت امکانی خصیصه مورد نظر در فرد i ام گوئیم. این نسبت در واقع، امکان داشتن خصیصه مورد نظر به عدم داشتن آن را برای هر فرد نشان می‌دهد. در مدل مورد بررسی، مشاهدات متغیر وابسته به صورت اعداد فازی μ_i و مشاهدات دقیق متغیرهای مستقل برای فرد i ام به صورت $j = 1, 2, \dots, m$; $i = 1, 2, \dots, n$ می‌باشند. مدل رگرسیون لوجستیک فازی به صورت زیر خواهد بود:

$$\hat{W}_i = \ln\left(\frac{\mu_i}{1-\mu_i}\right) = \sum_{j=0}^n A_j x_{ij}; \quad x_{i0} = 1; \quad i = 0, 1, \dots, m \quad (1)$$

که در آن $A_j = (l_j, c_j, r_j)$ ضرایب فازی مثلثی هستند. در این مدل‌ها ضرایب رگرسیونی لزوماً اعداد فازی مثلثی متقارن نیستند. l_j نقطه پایانی چپ، r_j نقطه

^{۱۱}Measure Of Performance based on fuzzy distance

پایانی راست و c_j مد A_j می باشند. با استفاده از اصل توسیع می توان برآورد خروجی ها $\hat{W}_i$ را مجدداً به صورت امکان موفقیت $\mu_i(x)$ تبدیل نمود. به طور مثال می توان درجات سازگاری را به صورت زیر (شکل ۱) در نظر گرفت:



شکل ۱: افراز پیشنهادی (μ_i) برای واژه های بیانی

$$\mu = (Low, Moderate, High)$$

$$Low = (0/01, 0/02, 0/3)$$

$$Moderate = (0/25, 0/5, 0/75)$$

$$High = (0/7, 0/98, 0/99) \quad (2)$$

در این مدل افراد خبره واژه های بیانی مانند (کم، متوسط، زیاد، ...) را به عنوان امکان موفقیت تعیین می کنند و تبدیل لگاریتمی بخت امکانی $i = \ln(\frac{\mu_i}{1-\mu_i})$ ، $i = 1, \dots, m$ به عنوان مقادیر مشاهده شده در نظر گرفته می شوند. تابع عضویت این مشاهدات خروجی با استفاده از تابع عضویت تعریف شده μ_i و اصل توسیع به صورت زیر محاسبه می شود:

$$w_i(y) = \sup_{\forall x: \ln \frac{x}{1-x} = y} \mu_i(x) \quad (3)$$

که در آن $0 < x < 1$ و $f(x) = \ln \frac{x}{1-x}$ تابعی یک به یک است. لذا در این حالت فقط و فقط یک $x \in (0, 1)$ وجود دارد که $\ln \frac{x}{1-x} = y$

$$w_i(y = \ln \frac{x}{1-x}) = \mu_i(\frac{\exp(x)}{1 + \exp(x)}) \quad (4)$$

برآورد ضرایب مدل نیز به دو روش کمترین قدر مطلق اختلافات^{۱۲} و کمترین مربعات^{۱۳} که در مطالعه (۱۰) استفاده شده، محاسبه می شود. معیار نیکویی برازش نیز شاخص قابلیت بر اساس فاصله های فازی است. نحوه ی برآورد ضرایب و محاسبه نیکویی برازش در پیوست آورده شده است.

^{۱۲}Least Absolute Deviations (LAD)

^{۱۳}Least Squares Estimation (LSE)

۳ کاربرد در مدلسازی شدت اختلال اوتیسم و نتایج

از ۲۰ کودک مبتلا به اختلال اوتیسم که در این پژوهش مورد مطالعه قرار گرفتند، ۳ نفر توانایی تکلم نداشته و یک نفر سوالات خرده مقیاس تعامل اجتماعی را پاسخ نداده بود، لذا این افراد از مطالعه خارج شدند. میانگین سنی ۱۶ کودک باقی مانده در مطالعه ۶۳/۱۰ سال با انحراف معیار ۴۸/۳ می باشد. چهار کودک دختر (۲۵ درصد) و دوازده نفر دیگر (۷۵ درصد) پسر هستند. نمره ی خام خرده مقیاس های حرکات کلیشه ای، ارتباط و تعامل اجتماعی به عنوان سه متغیر ورودی و نظر درمانگر در مورد شدت اختلال فرد که به صورت واژه های بیانی (خفیف، متوسط و شدید) می باشد، به عنوان متغیر خروجی در مدل وارد شده اند. افزاز پیشنهادی برای واژه های بیانی را به صورت شکل ۱ در نظر گرفته و اطلاعات جمع آوری شده در جدول ۱ آورده شده است. مدل برازش داده حاصل از برآورد ضرایب به روش های LAD و LSE، که در آن نمره خرده مقیاس حرکات کلیشه ای x_{i1} ، ارتباط x_{i2} و تعامل اجتماعی x_{i3} می باشد:

روش کمترین قدرمطلق اختلافات LSE :

$$\hat{W}_i = (-4/317, -3/099, -2/966) + (0/047, 0/047, 0/047)x_{i1} + (0/009, 0/009, 0/0121)x_{i2} + (0/142, 0/142, 0/217)x_{i3}$$

روش کمترین مربعات LAD :

$$\hat{W}_i = (-7/169, -5/626, -5/619) + (0/117, 0/117, 0/145)x_{i1} + (0/013, 0/013, 0/013)x_{i2} + (0/153, 0/153, 0/18)x_{i3}$$

ضریب ۰/۱۱۷ برای خرده مقیاس حرکات کلیشه ای (در مدل برآورد شده به روش LAD) نشان دهنده این است که با افزایش نمره این خرده مقیاس، بخت شدت اختلال اوتیسم افزایش می یابد. یعنی با فرض ثابت ماندن مقدار سایر خرده مقیاس ها، با هر واحد افزایش در نمره خرده مقیاس حرکات کلیشه ای تقریباً ۰/۱۱۷ یخت امکانی شدت اختلال افزایش می یابد.

جدول ۱: مشاهدات فازی شدت اختلال اوتیسم و نمرات خرده مقیاس ها.

ردیف	حرکات کلیشه ای	ارتباط	تعامل اجتماعی	نظر درمانگر
۱	۱۸	۴	۳۱	شدید
۲	۱۸	۱۵	۲۲	متوسط
۳	۲۱	۲۳	۲۲	متوسط
۴	۱۸	۳۵	۲۵	شدید
۵	۲۷	۲۳	۲۲	متوسط
۶	۷	۳	۱۱	متوسط
۷	۲۴	۲۸	۱۶	متوسط
۸	۲۲	۳۸	۱۹	متوسط
۹	۱۰	۱۲	۱۳	خفیف
۱۰	۱۲	۱۸	۲۰	خفیف
۱۱	۲۱	۲۳	۲۵	شدید
۱۲	۱۴	۲۴	۳۱	شدید
۱۳	۳۵	۲۸	۳۳	شدید
۱۴	۲۴	۱۴	۸	متوسط
۱۵	۸	۸	۹	متوسط
۱۶	۱۰	۹	۳	متوسط

معیار نیکویی برازش نیز نشان می دهد مدل برآورد شده به روش LSE بر داده های این مطالعه نسبت به روش LAD مناسب تر است. نتایج را می توان در جدول

۲ مشاهده نمود.

با استفاده از مدل های برآورد شده می توان لگاریتم بخت امکانی را برای مشاهده جدید نیز محاسبه نمود. مثلا برای فردی با :

نمره خرده مقیاس حرکات کلیشه ای = ۲۰

نمره خرده مقیاس ارتباط = ۱۲

نمره خرده مقیاس تعامل اجتماعی = ۳۰

لگاریتم بخت امکانی اختلال شدید برای این فرد به روش LAD (۰/۰۸۳, ۱/۴۶, ۲/۸۳۷) برآورد می شود. لذا، بخت امکانی اختلال شدید برای این فرد

تقریبا ۴/۳۱ می باشد. با استفاده از اصل توسیع تابع عضویت بخت امکانی شدت اختلال برای این فرد به صورت زیر خواهد بود:

$$\exp(\hat{W}_{New}(x)) = \frac{\mu_{New}}{1 - \mu_{New}}(x) = \begin{cases} 1 - \frac{1/46 - \ln(x)}{1/543} & 0/92 \leq x \leq 4/31 \\ 1 - \frac{\ln(x) - 1/46}{1/377} & 4/31 < x \leq 17/06 \end{cases}$$

$$\hat{\mu}_{New}(x) = \hat{W}_{New}\left(\ln \frac{x}{1-x}\right) = \begin{cases} 1 - \frac{1/46 - \ln(\frac{x}{1-x})}{1/543} & 0/479 \leq x \leq 0/812 \\ 1 - \frac{\ln(\frac{x}{1-x}) - 1/46}{1/377} & 0/812 < x \leq 0/945 \end{cases}$$

جدول ۲: برآورد ضرایب و معیار نیکویی برازش.

ضرایب	روش LAD	روش LSE
(l_0, c_0, r_0)	$(-7/169, -5/626, -5/619)$	$(-4/317, -3/099, -2/966)$
(l_1, c_1, r_1)	$(0/117, 0/117, 0/145)$	$(0/047, 0/047, 0/047)$
(l_2, c_2, r_2)	$(0/013, 0/013, 0/013)$	$(0/009, 0/009, 0/0121)$
(l_3, c_3, r_3)	$(0/153, 0/153, 0/18)$	$(0/142, 0/142, 0/217)$
M_p	0/2310	0/0920

اگر بتوان نظر درمانگر در خصوص شدت بیماری را به صورت اطلاعات دقیق در نظر گرفت، می توان از رگرسیون لجستیک برای مدل سازی استفاده نمود.

سپاسگزاری

بدینوسیله از مربیان و درمانگران پیش دبستانی و دبستان غیردولتی اختلال رفتاری - هیجانی و طیف اوتیسم پرورش گرگان و همچنین خانواده ی کودکانی که از اطلاعات آنها در این مطالعه استفاده شده است، به علت همکاری و صبوری در تکمیل پرسشنامه ها تشکر می شود.

پیوست : مروری کوتاه بر رگرسیون امکانی

برآورد ضرایب مدل:

به منظور کمینه کردن اختلاف بین مشاهدات و پیشبینی های فازی، که هدف رگرسیون فازی است، می توان از α - برش های خروجی های فازی مشاهده شده و

پیشبینی شده استفاده نمود.

α -برش های تبدیل لگاریتمی بخت مشاهدات $W_i(\alpha)$ را می توان محاسبه نمود. به طور مثال تابع عضویت $\mu_i = (l_{\mu_i}, m_{\mu_i}, r_{\mu_i})$ را برای بخت امکانی داشتن خصیصه مورد نظر برای فرد i ام را در نظر بگیریم. آنگاه α -برش ها را می توان به صورت زیر به دست آورد:

$$\mu_i(\alpha) = [l_{\mu_i}, r_{\mu_i}] = [m_{\mu_i} - (\alpha)(m_{\mu_i} - l_{\mu_i}), m_{\mu_i} + (\alpha)(r_{\mu_i} - m_{\mu_i})] \quad (5)$$

به منظور دستیابی به $W_i(\alpha)$ ، داریم:

$$W_i(\alpha) = [l_{w_i}, r_{w_i}] = [-\ln(\frac{1-l_{\mu_i}}{l_{\mu_i}}), -\ln(\frac{1-r_{\mu_i}}{r_{\mu_i}})] \quad (6)$$

فاصله زیر نشان دهنده α - برش های تبدیل لگاریتمی بخت های برآورد شده $\hat{W}_i(\alpha)$ است:

$$\hat{W}_i(\alpha) = [l_{\hat{w}_i}, r_{\hat{w}_i}] = \sum_{j=0}^p [l_{A_j}(\alpha), r_{A_j}(\alpha)] x_{ij} = [\sum_{j=0}^p l_{A_j}(\alpha) x_{ij}, \sum_{j=0}^p r_{A_j}(\alpha) x_{ij}] \quad (7)$$

روش کمترین قدرمطلق اختلافات:

برآورد های کمترین قدر مطلق اختلافات $\bar{l}_{A_j}(\alpha)$ و $\bar{r}_{A_j}(\alpha)$ را به طوری که در شرایط زیر صدق کنند، می توان محاسبه نمود:

$$\begin{aligned} \sum_{i=1}^n |l_{W_j}(\alpha) - \sum_{j=0}^p l_{A_j}(\alpha) x_{ij}| &= \min! \\ \sum_{i=1}^n |r_{W_j}(\alpha) - \sum_{j=0}^p r_{A_j}(\alpha) x_{ij}| &= \min! \end{aligned} \quad (8)$$

روش کمترین مربعات:

برای برآورد ضرایب به روش کمترین مربعات، فرمول های روش کمترین قدرمطلق اختلافات با فرمول های زیر جایگزین می شوند:

$$\begin{aligned} \sum_{i=1}^n (l_{W_j}(\alpha) - \sum_{j=0}^p l_{A_j}(\alpha) x_{ij})^2 &= \min! \\ \sum_{i=1}^n (r_{W_j}(\alpha) - \sum_{j=0}^p r_{A_j}(\alpha) x_{ij})^2 &= \min! \end{aligned} \quad (9)$$

حل مسائل کمینه سازی با استفاده از دستور NMinimize در نرم افزار Mathematica 10.4.0 انجام گرفته است.

معیار نیکویی برازش M_P :

در تحلیل رگرسیون کلاسیک، میانگین قدرمطلق خطاها به عنوان روشی برای ارزیابی نیکویی برازش مورد استفاده قرار می گیرد. در محیط فازی، بر اساس اختلاف بین مقادیر عضویت اعداد فازی مشاهده شده و اعداد فازی برآورد شده، میتوانیم از میانگین قدرمطلق درصد خطاها به عنوان معیاری برای نیکویی برازش استفاده نماییم. اما زمانی که مقادیر مشاهده شده و برآورد شده همپوشانی ندارند، کارایی این معیار کم است. در این حالت، شاخص قابلیت نیز بدون در نظر گرفتن فاصله بین مقادیر مشاهده شده و برآورد شده، کارایی مناسبی ندارد. برای حل این مشکل، از شاخص اصلاح شده زیر به نام شاخص قابلیت بر اساس فاصله های فازی استفاده می کنیم. هر چه اختلاف بین دو عدد فازی کمتر باشد، مقدار M_P به صفر نزدیک تر خواهد بود، که یعنی دقت مدل بالاتر است (۱۰).

$$M_P = \frac{1}{n} \sum_{i=0}^n d(W_i, \hat{W}_i) \quad (10)$$

که در آن: $d(W_i, \hat{W}_i) = \frac{\int |\mu_{w_i}(x) - \mu_{\hat{w}_i}(x)| dx}{\int \mu_{w_i}(x) dx} + h_d(W_i(\circ), \hat{W}_i(\circ))$ و $h_d(A, B) = \inf_{a \in A} \inf_{b \in B} |a - b|$

- [۱] احمدی، صفری، همتیان و ... (۲۰۱۱). بررسی شاخص های روان سنجی آزمون تشخیصی اوتیسم (GARS)، پژوهش های علوم شناختی و رفتاری، ۱(۱)، ۸۷-۱۰۴
- [۲] زاده، عابدی و & اژی (۲۰۱۷). بررسی روایی و پایایی نسخه فارسی پرسش نامه غربالگری زود هنگام ویژگی های اوتیسم (ESAT-PV) در کودکان نوپا، فصلنامه علمی پژوهشی توانبخشی، ۱۸(۳)، ۱۹۳-۱۸۲
- [3] Association, A. P. (2013), Diagnostic and Statistical Manual of Mental Disorders (DSM-5®): *American Psychiatric Publishing*.
- [4] Baio, J. (2012), Prevalence of Autism Spectrum Disorders: Autism and Developmental Disabilities Monitoring Network, 14 Sites, United States, 2008. Morbidity and Mortality Weekly Report, *Surveillance Summaries. Volume 61, Number 3. Centers for Disease Control and Prevention*.
- [5] Baio, J. (2014), Prevalence of autism spectrum disorder among children aged 8 years-autism and developmental disabilities monitoring network, 11 sites, United States, 2010, *Surveillance Summaries. Centers for Disease Control and Prevention*.
- [6] Bozorgnia, A., Malekpour, M., & Abedi, A. (2011), Prevalence of autism in children 6 to 12 years old Shhrkrd 2009-2010, *Paper presented at the Regional Conference on Child and Adolescent Psychology*.
- [7] Ganji, M. (2013), Abnormal psychology based on DSM-5, *Tehran: savalan Publications*, 260-261.
- [8] Kanner, L. (1943), Autistic disturbances of affective contact, *Nervous child*, 2(3), 217-250.
- [9] Namdari, M., Taheri, S. M., Abadi, A., Rezaei, M., & Kalantari, N. (2013), Possibilistic logistic regression for fuzzy categorical response data, *Paper presented at the Fuzzy Systems (FUZZ), IEEE International Conference on 2013*.
- [10] Namdari, M., Yoon, J. H., Abadi, A., Taheri, S. M., & Choi, S. H. (2015), Fuzzy logistic regression with least absolute deviations estimators, *Soft Computing*, 19(4), 909-917.
- [11] Pourahmad, S., Ayatollahi, S. M. T., Taheri, S. M., & Agahi, Z. H. (2011), Fuzzy logistic regression based on the least squares approach with application in clinical studies, *Computers & Mathematics with Applications*, 62(9), 3353-3365.
- [12] Pourahmad, S., Ayatollahi, T., Mohammad, S., & Taheri, S. M. (2011), Fuzzy logistic regression: a new possibilistic model and its application in clinical vague status, *Iranian Journal of Fuzzy Systems*, 8(1), 1-17.
- [13] Rice, C. (2009), Prevalence of autism spectrum disorders-Autism and developmental disabilities monitoring network, United States, 2006.
- [14] Samadi, S. A., Mahmoodizadeh, A., & McConkey, R. (2012), A national study of the prevalence of autism among five-year-old children in Iran, *Autism*, 16(1), 5-14.

- [15] Samadi, S. A., & McConkey, R. (2015), Screening for autism in Iranian preschoolers: Contrasting M-CHAT and a scale developed in Iran, *Journal of Autism and developmental disorders*, 45(9), 2908-2916.
- [16] SARABI, J. M., Hasanadadi, A., Mashhadi, A., & Asgharinekah, M. (2012), The effects of parent education and skill training program on stress of mothers of children with autism.
- [17] Soto, S., Linas, K., Jacobstein, D., Biel, M., Migdal, T., & Anthony, B. J. (2015), A review of cultural adaptations of screening tools for autism spectrum disorders, *Autism*, 19(6), 646-661.
- [18] Volkmar, F. R., Paul, R., Klin, A., & Cohen, D. J. (2007), Handbook of Autism and Pervasive Developmental Disorders, Assessment, Interventions, and Policy, *Wiley*.
- [19] Zadeh, L. A. (1969), Biological application of the theory of fuzzy sets and systems, *Paper presented at the The Proceedings of an International Symposium on Biocybernetics of the Central Nervous System*.
- [20] Zadeh, L. A. (1973), Outline of a new approach to the analysis of complex systems and decision processes, *IEEE Transactions on systems, Man, and Cybernetics*(1), 162(2), 28-44.

برآورد پارامتر توزیع نمایی به کمک الگوریتم EM مبتنی بر مشاهدات فازی

عباس پرچی^۱

دانشگاه شهید باهنر کرمان، دانشکده ریاضی و کامپیوتر، بخش آمار

چکیده

الگوریتم EM ابزاری قدرتمند برای برآورد بیشترین درستنمایی مبتنی بر داده‌های ناقص است که در اغلب کتاب‌های استنباط آماری نیز مطرح شده است. در اینجا معنی کلمه «ناقص» حالتی کلی دارد و در موقعیت‌های متفاوت می‌تواند به معانی گوناگونی (مانند داده‌های گم شده، داده‌های بازه‌ای، مشاهدات سانسور شده و نظایر آنها) اشاره داشته باشد. این مقاله کاربردی جدیدی از الگوریتم EM را مطرح می‌سازد که در آن منظور از داده‌های ناقص داده‌های نادقیق/فازی هستند. بر اساس این نوع داده‌های مبهم، در این مقاله برآورد بیشینه درستنمایی پارامتر توزیع نمایی به کمک الگوریتم EM در قالب یک مثال عددی محاسبه شده است و این مثال می‌تواند به منظور فهم مطلب و نیز استفاده از الگوریتم EM در مثال‌ها/حالت‌های پیچیده‌تر، برای دانشجویان تحصیلات تکمیلی مفید باشد.

کلمات کلیدی: الگوریتم EM ، برآورد بیشینه درستنمایی، داده‌های فازی.

۱ مقدمه‌ای بر الگوریتم EM

فرض کنید $\mathbf{Y}$ بردار داده‌های مشاهده شده و $\mathbf{X}$ بردار داده‌های فازی باشند. همچنین فرض کنید θ پارامتر مجهول است که می‌خواهیم برآورد شود و $l_c(\theta; \mathbf{Y}, \mathbf{X})$ لگاریتم تابع درستنمایی بر اساس کلیه داده‌ها است که به ازای تمام مقادیر ممکن θ در فضای پارامتر Ω ، تعریف می‌شود. الگوریتم EM با یک مقدار اولیه $\theta^{(0)} \in \Omega$ آغاز می‌شود و دو گام زیر را تا رسیدن به همگرایی، تکرار می‌کند:

گام E ^۱: محاسبه $l_c^{(j)}(\theta) = E_{\mathbf{X}|\mathbf{Y}, \theta^{(j-1)}} [l_c(\theta; \mathbf{Y}, \mathbf{X})]$ که امید ریاضی با توجه به داده‌های ناقص $\mathbf{X}$ به شرط داده‌های مشاهده شده $\mathbf{Y}$ گرفته می‌شود و باید توجه شود که مقدار $\theta^{(j-1)}$ در این امید ریاضی جایگذاری می‌شود.

^۱عباس پرچی : parchami@uk.ac.ir

^۱Expectation

گام ۲: یافتن $\theta^{(j)} \in \Omega$ بقسمی که $l_c^{(j)}(\theta)$ بیشینه شود.

تکرار دو گام فوق به ازای $j = 1, 2, \dots$ ، منجر به همگرایی دنباله $\theta^{(1)}, \theta^{(2)}, \dots$ در ماکزیم موضعی لگاریتم درستنمایی کلیه داده‌های کامل می‌شود. همان‌گونه که مطرح شد، الگوریتم EM ابزاری برای برآورد بیشینه درستنمایی مبتنی بر داده‌های ناقص است و داده‌های فازی می‌تواند یکی از مصداق‌های داده‌های ناقص به حساب آید، زیرا در این حالت بجای اینکه مقدار دقیقی به عنوان مشاهده به ثبت رسیده باشد تنها یک مقدار تقریبی (در قالب یک عدد فازی) مشاهده و ثبت شده است.

مثال‌های اندکی از این وجه کاربردی الگوریتم EM که نخستین بار توسط (۱) پیشنهاد شد، در دست است و در بخش ۳ قصد داریم تا از این رویکرد در برآورد بیشترین درستنمایی پارامتر مجهول توزیع نمایی با استفاده از الگوریتم EM استفاده کنیم. لذا برخی از تعاریف و مفاهیم پیش‌نیاز را در بخش ۲ مطرح می‌سازیم.

۲ احتمال شرطی بر اساس مشاهدات فازی

در این بخش قصد داریم تا برخی از مفاهیم و تعاریف مقدماتی که در بخش بعد مورد نیاز است را مطرح سازیم.

تعریف ۱.۲. تابع چگالی/جرم احتمال شرطی X به شرط مشاهده‌ی فازی $\tilde{x}$ را با نماد $f_{\theta}(x|X \in \tilde{x})$ نمایش و بصورت زیر تعریف می‌کنیم

$$f_{\theta}(x|X \in \tilde{x}) = \frac{\tilde{x}(x) f_{\theta}(x)}{\int \tilde{x}(x) f_{\theta}(x) dx} \quad (۱)$$

در حالت گسسته به جای انتگرال از مجموع‌یابی استفاده می‌شود. لازم بذکر است که (۳) تعریف مشابهی تحت عنوان چگالی X بر اساس اطلاعات فازی مطرح کرده است.

توجه ۲.۲. بدیهی است که تابع چگالی/جرم احتمال شرطی معرفی شده در تعریف ۱.۲، از خواص یک تابع چگالی/جرم احتمال مانند $f_{\theta}(x|X \in \tilde{x}) \geq 0, \forall x \in R$ و $\int_{-\infty}^{\infty} f_{\theta}(x|X \in \tilde{x}) dx = 1$ برخوردار است.

نتیجه ۳.۲. فرض کنید $\mathbf{X} = (X_1, \dots, X_n)$ نمونه‌ای تصادفی از یک جامعه با تابع چگالی/جرم احتمال $f_{\theta}(x)$ و همچنین $\tilde{\mathbf{X}} = (\tilde{x}_1, \dots, \tilde{x}_n)$ مقدار مشاهده‌شده‌ی نمونه باشد، بطوریکه $\tilde{x}_i$ مقدار مشاهده‌شده‌ی فازی برای متغیر X_i با تابع عضویت $\tilde{x}(x)$ به ازای $i = 1, \dots, n$ است. بعنوان تابعی از θ ، تابع درستنمایی مبتنی بر مشاهدات فازی $\tilde{\mathbf{X}}$ برابر است با

$$\begin{aligned} L(\theta|\tilde{\mathbf{X}}) &= L(\theta|\tilde{x}_1, \dots, \tilde{x}_n) \\ &= \prod_{i=1}^n P_{\theta}(X_i \in \tilde{x}_i) \\ &= \prod_{i=1}^n \int_{-\infty}^{\infty} \tilde{x}_i(x) f_{\theta}(x) dx \end{aligned} \quad (۲)$$

نتیجه ۴.۲. احتمال شرطی پیشامد فازی $\tilde{x}'$ ، بر اساس احتمال پیشامد فازی (۴)، معادل است با

$$\begin{aligned} P_{\theta}(X \in \tilde{x}'|X \in \tilde{x}) &= \frac{\int \tilde{x}'(x) f_{\theta}(x|X \in \tilde{x}) dx}{\int \tilde{x}'(x) \tilde{x}(x) f_{\theta}(x) dx} \\ &= \frac{\int \tilde{x}'(x) f_{\theta}(x) dx}{\int \tilde{x}(x) f_{\theta}(x) dx} \\ &= \frac{P_{\theta}(X \in (\tilde{x}' \cap \tilde{x}))}{P_{\theta}(X \in \tilde{x})} \end{aligned} \quad (۳)$$

که در آن تابع عضویت اشتراک دو مجموعه فازی $\tilde{x}'$ و $\tilde{x}$ بصورت $(\tilde{x}' \cap \tilde{x})(x) = \tilde{x}'(x) \tilde{x}(x)$ تعریف شده است.

^۲Maximization

نتیجه ۵.۲. بر اساس تابع چگالی/جرم احتمال شرطی معرفی شده در تعریف ۱.۲، امید ریاضی متغیر تصادفی X به شرط مشاهده‌ی فازی $\tilde{x}$ بصورت زیر بدست می‌آید (در حالت گسسته به جای انتگرال از مجموعیابی استفاده شود)

$$\begin{aligned} E_{\theta}(X | X \in \tilde{x}) &= \int_{-\infty}^{\infty} x f_{\theta}(x | X \in \tilde{x}) dx \\ &= \frac{\int x \tilde{x}(x) f_{\theta}(x) dx}{\int \tilde{x}(x) f_{\theta}(x) dx} \end{aligned} \quad (۴)$$

۳ تشریح یک مثال عددی برای الگوریتم EM مبتنی بر داده‌های ناقص فازی-مقدار

فرض کنید $X = (X_1, \dots, X_n)$ یک نمونه تصادفی از توزیع نمایی با پارامتر λ باشد. حالتی را در نظر بگیرید که تمامی مشاهدات از نوع فازی بوده و در نتیجه بجای مشاهده‌ی مقادیر دقیق x_i ها، یک تابع عضویت برای هر داده‌ی فازی به ثبت رسیده است. به عبارت دیگر، در این مثال قصد داریم تا برآورد بیشینه درستنمایی پارامتر مجهول λ را مبتنی بر مشاهدات/داده‌های فازی-مقدار $\tilde{\mathbf{X}} = (\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_n)$ به کمک الگوریتم EM محاسبه کنیم و لذا طبق الگوریتم مطرح شده در بخش ۱، بدیهی است که $\mathbf{X} = \tilde{\mathbf{X}}$ و $\mathbf{Y} = \phi$. تابع درستنمایی مبتنی بر این داده‌های فازی-مقدار برابر است با

$$L(\lambda | \tilde{\mathbf{X}}) = \prod_{i=1}^n P_{\lambda}(X_i \in \tilde{x}_i) = \prod_{i=1}^n \int_0^{\infty} \tilde{x}_i(x) f_{\lambda}(x) dx \quad (۵)$$

گرچه برای یافتن $MLE(\lambda)$ نمی‌توان از این تابع (و یا لگاریتم آن) نسبت به λ مشتق گرفت، اما می‌توان برای بیشینه‌سازی آن از روش‌های عددی (هم‌چون روش نیوتن رافسون) استفاده کرد. به منظور استفاده از الگوریتم EM ، که رویکرد مورد نظر در این مقاله است، ابتدا تابع لگاریتم درستنمایی را مبتنی بر داده‌های کامل/حقیقی-مقدار محاسبه می‌کنیم

$$L_c(\lambda) = \lambda^n \exp(-\lambda \sum_{i=1}^n x_i) \Rightarrow l_c(\lambda) = n \ln(\lambda) - \lambda \sum_{i=1}^n x_i \quad (۶)$$

و با برابر صفر قرار دادن مشتق لگاریتم درستنمایی، $MLE(\lambda) = \frac{1}{\bar{x}} = \frac{n}{\sum_{i=1}^n x_i}$ بدست می‌آید. اما به یاد داشته باشیم که هدف محاسبه‌ی برآورد بیشینه درستنمایی بر اساس داده‌های فازی (و نه حقیقی) است.

طبق تعریف ۱.۲، تابع چگالی احتمال شرطی X_i به شرط مشاهده‌ی فازی $\tilde{x}_i$ برابر است با

$$\begin{aligned} f_{\lambda}(x_i | X_i \in \tilde{x}_i) &= \frac{\tilde{x}_i(x_i) f_{\lambda}(x_i)}{\int \tilde{x}_i(x) f_{\lambda}(x) dx_i} \\ &= \frac{\tilde{x}_i(x_i) \lambda \exp(-\lambda x_i)}{\int \tilde{x}_i(x) \lambda \exp(-\lambda x) dx} \\ &= \frac{\tilde{x}_i(x_i) \exp(-\lambda x_i)}{\int \tilde{x}_i(x) \exp(-\lambda x) dx}, \quad i = 1, \dots, n. \end{aligned} \quad (۷)$$

لگاریتم درستنمایی J امین مرحله از الگوریتم مبتنی بر مشاهدات فازی (گام E) برابر

$$\begin{aligned} l_c^{(j)}(\lambda) &= E_{\lambda^{(j-1)}} [l_c(\lambda) | X_1 \in \tilde{x}_1, \dots, X_n \in \tilde{x}_n] \\ &= n \ln(\lambda) - \lambda \sum_{i=1}^n E_{\lambda^{(j-1)}}(X_i | X_i \in \tilde{x}_i) \end{aligned} \quad (۸)$$

است و مقدار بیشینه‌ی آن (گام M)، با ثابت فرض کردن $\lambda^{(j-1)}$ (چون مقدارش در مرحله قبلی الگوریتم تعیین شده)، معادل است با

$$\frac{\partial l_c^{(j)}(\lambda)}{\partial \lambda} = 0 \Rightarrow \hat{\lambda}^{(j)} = \frac{n}{\sum_{i=1}^n E_{\lambda^{(j-1)}}(X_i | X_i \in \tilde{x}_i)} = \frac{n}{\sum_{i=1}^n \hat{x}_i^{(j-1)}} \quad (۹)$$

که در آن، مطابق نتیجه ۵.۲، داریم

$$\begin{aligned}
 \hat{x}_i^{(j-1)} &= E_{\lambda^{(j-1)}}(X_i | X_i \in \tilde{x}_i) \\
 &= \int_{-\infty}^{\infty} x f_{\lambda^{(j-1)}}(x | X_i \in \tilde{x}_i) dx \\
 &= \frac{\int_{-\infty}^{\infty} x \tilde{x}_i(x) f_{\lambda^{(j-1)}}(x) dx}{\int_{-\infty}^{\infty} \tilde{x}_i(x) f_{\lambda^{(j-1)}}(x) dx} \\
 &= \left\{ \frac{\int x \tilde{x}_i(x) \exp(-\lambda x) dx}{\int \tilde{x}_i(x) \exp(-\lambda x) dx} \right\}_{\lambda=\lambda^{(j-1)}} \quad (10)
 \end{aligned}$$

اگر چه روش مطرح شده در این مقاله کلیت داشته و برای هر نوع عدد فازی برقرار است، اما به دلیل سادگی مطالب در ادامه از داده‌های مثلثی فازی با نماد $\tilde{x} = T(n, l_n, r_n)$ و تابع عضویت زیر استفاده شده است که در آن n ، l_n و r_n به ترتیب نشان دهنده‌ی هسته، پهنای چپ و پهنای راست عدد فازی $\tilde{x}$ هستند

$$\tilde{x}(x) = \begin{cases} \frac{x-n+l_n}{l_n} & \text{اگر } n-l_n \leq x < n \\ \frac{n+r_n-x}{r_n} & \text{اگر } n \leq x < n+r_n \\ 0 & \text{جاهای دیگر} \end{cases} \quad (11)$$

نمونه‌ای تصادفی از داده‌های مثلثی فازی زیر را برای برآورد پارامتر مجهول λ در نظر بگیرید

$$\tilde{x}_1 = T(1/0.19, 0/0.865, 1/0.374)$$

$$\tilde{x}_2 = T(1/0.141, 0/0.316, 1/0.548)$$

$$\tilde{x}_3 = T(0/0.806, 0/0.523, 0/0.944)$$

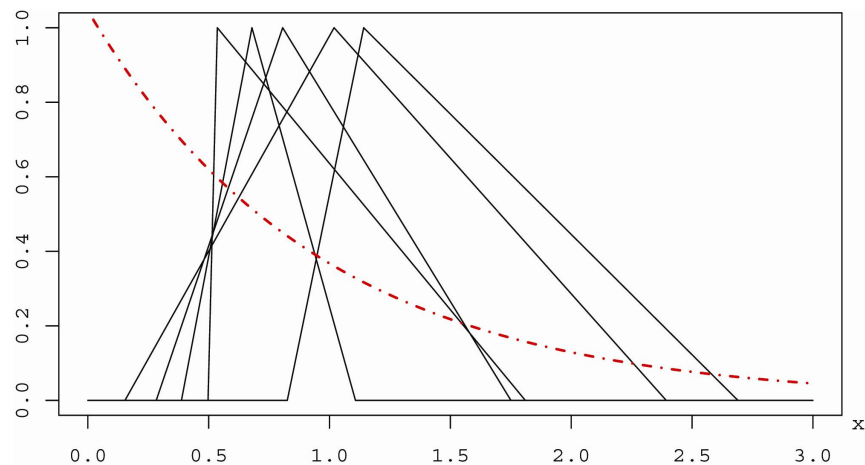
$$\tilde{x}_4 = T(0/0.679, 0/0.292, 0/0.428)$$

$$\tilde{x}_5 = T(0/0.536, 0/0.038, 1/0.274)$$

جدول ۱ تکرار الگوریتم و رسیدن به همگرایی آن را مبتنی بر این داده‌های فازی، به ازای دو مقدار آغازین متفاوت $\lambda_1^{(0)}$ و $\lambda_3^{(0)}$ نشان می‌دهد که بر اساس هر دو مقدار آغازین، الگوریتم EM به یک همگرایی ($\hat{\lambda} = 1/0.42$) رسیده است.

جدول ۱: روند همگرایی الگوریتم EM به ازای دو مقدار آغازین مختلف در مثال ۱.

$\lambda_3^{(j)}$	$\lambda_1^{(j)}$	شماره تکرار الگوریتم (j)
۳۰	۵	۰
۲/۰.۳۱	۱/۰.۳۹۷	۱
۱/۰.۴۳	۱/۰.۷۹	۲
۱/۰.۵۲	۱/۰.۴۵	۳
۱/۰.۴۳	۱/۰.۴۲	۴
۱/۰.۴۲	۱/۰.۴۲	۵
۱/۰.۴۲		۶



شکل ۱: توابع عضویت داده‌های مثلثی فازی به همراه تابع چگالی احتمالی نمایشی برآورد شده.

مراجع

- [1] Denoeux, T. (2011), Maximum likelihood estimation from fuzzy data using the EM algorithm, *Fuzzy Sets and Systems* 183, 72–91.
- [2] Knight, K. (2000), *Mathematical Statistics*, New York, Chapman & Hall.
- [3] Pourmousa, R. (2018), On truncated measures of income inequality from a fuzzy perspective, *Iranian Journal of Fuzzy Systems* 15, 123-137.
- [4] Zadeh, L.A. (1968), Probability measures of fuzzy events, *J. Math. Anal. Appl.* 23, 421–427.

پیشنهادی برای رفع چالش متغیر تصادفی فازی

عباس پرچمی^۱

دانشگاه شهید باهنر کرمان، دانشکده ریاضی و کامپیوتر، بخش آمار

چکیده

پس از مروری انتقادی به متغیرهای تصادفی فازی که تا کنون مطرح شده‌اند، در این مقاله نوعی جدید از این مفهوم به نام «متغیر تصادفی قطعه‌قطعه‌شده‌ی خطی» مبتنی بر تعدادی متناهی از برش‌های تصادفی معرفی شده است. وقتی تعداد این برش‌های تصادفی به بینهایت میل می‌کند، «متغیر تصادفی قطعه‌قطعه‌شده‌ی خطی» به نوعی دیگر از متغیر تصادفی فازی، که آن را «متغیر تصادفی فازی LR » می‌نامیم، میل می‌کند. بنابراین مبتنی بر این حالت حدی، در انتهای این مقاله تعریفی بهتر، دقیق‌تر و در عین حال ساده‌تر برای متغیر تصادفی فازی ارائه می‌شود. همچنین، چندین مثال عددی برای انتقال بهتر مفاهیم و تعاریف، بر اساس شبیه‌سازی مطرح شده است.

کلمات کلیدی: متغیر تصادفی فازی، عدد فازی LR ، عدد فازی قطعه‌قطعه‌شده‌ی خطی، تابع توزیع تجربی پیوسته.

۱ اهم مطالعات انجام‌شده در خصوص متغیر تصادفی فازی

تا کنون چندین رویکرد مختلف برای مدل متغیرهای تصادفی فازی پیشنهاد شده است که در این بخش ایده‌ی برخی از آنها را مطرح و مورد انتقاد قرار می‌دهیم.

۱.۱ متغیر تصادفی فازی مبتنی بر برش‌های تصادفی

در این رویکرد، تابع فازی-مقدار $\tilde{X} : \Omega \rightarrow F(\mathbb{R})$ یک متغیر تصادفی فازی در فضای احتمال (Ω, A, P) نامیده می‌شود هرگاه اندازه پذیر باشد. ثابت شده است که $\tilde{X} : \Omega \rightarrow F(\mathbb{R})$ یک متغیر تصادفی فازی است اگر و تنها اگر برای هر $\alpha \in [0, 1]$ ، توابع $\tilde{X}_\alpha^l : \Omega \rightarrow \mathbb{R}$ و $\tilde{X}_\alpha^r : \Omega \rightarrow \mathbb{R}$ متغیرهای تصادفی حقیقی-مقدار باشند (که در آن $(\forall \omega \in \Omega; \tilde{X}(\omega)_\alpha = [\tilde{X}_\alpha^l(\omega), \tilde{X}_\alpha^r(\omega)])$). (۷:۸)

^۱عباس پرچمی : parchami@uk.ac.ir

مشکل این ایده آنست که احتمال یک مشاهده‌ی فازی مبتنی بر این تعریف از متغیر تصادفی فازی برابر صفر است. به عبارت دقیق‌تر، در این روش احتمال تودرتو بودن برش‌های تولید/شبیه‌سازی شده از یک متغیر تصادفی فازی (که شرطی لازم برای یک مجموعه فازی است) برابر صفر می‌باشد. همچنین تضمینی نیست که پس از شبیه‌سازی نقاط ابتدایی ($\tilde{X}_\alpha^l$) و انتهایی ($\tilde{X}_\alpha^r$) برش دلخواه α ، رابطه $\tilde{X}_\alpha^l \leq \tilde{X}_\alpha^r$ (که شرط لازم برای بازه بودن $[\tilde{X}_\alpha^l, \tilde{X}_\alpha^r]$ است) برقرار باشد.

۲.۱ متغیر تصادفی با پارامتر فازی-مقدار

بعضی از نویسندگان، مانند (۵؛ ۶)، متغیر تصادفی فازی را به‌گونه‌ای تعریف کردند که تصادفی بودن این متغیر فازی-مقدار ریشه در تنها یک متغیر تصادفی معمولی داشت. مثلاً به اعتقاد آن‌ها، $\tilde{X}$ یک متغیر تصادفی فازی نرمال با میانگین $\tilde{\theta}$ و واریانس ۱ است، اگر و تنها اگر $\tilde{X} = \tilde{\theta} + \xi$ که در آن ξ یک متغیر تصادفی غیرفازی نرمال استاندارد می‌باشد. گرچه مشاهدات مربوط به این گونه متغیر تصادفی فازی، بصورت فازی است، اما ابهام و نیز توابع شکل برای تمامی مشاهدات فازی-مقدار ثابت بوده و ماهیت تصادفی ندارد.

۳.۱ مشاهده‌ی فازی از یک متغیر تصادفی دقیق

(۴) برای تعمیم آمارها مبتنی بر داده‌های فازی، فرض کردند که متغیر تصادفی دقیق است و از آنجا که ابزاری دقیق برای سنجش/اندازه‌گیری مقدار آن در دست نداشتند، مقدار مشاهده‌ی این متغیر تصادفی دقیق را یک عدد نادقیق/فازی در نظر گرفتند. مجدداً تاکید می‌کنیم که در این حالت متغیر تصادفی بصورت فازی تعریف نشده است و لذا از مصداق‌های کاربردی کمی برخوردار است.

۲ متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی

در این بخش قصد داریم تا نوعی جدید از متغیر تصادفی فازی، با نام متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی را تعریف کنیم که در بخش‌های بعدی این مقاله زیربنای تعریف متغیر تصادفی فازی LR می‌شود. متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی در حقیقت یک چند-بر/چند-ضلعی متشکل از اتصال چندین پاره‌خط است که مراحل شبیه‌سازی آن در تعریف زیر ارائه شده است.

تعریف ۱.۲. مراحل تولید/شبیه‌سازی متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی

$$\tilde{X}_{knot.n} \approx [C_X \sim f_C, S_X^l \sim f_{S_X^l}, S_X^r \sim f_{S_X^r}]_{knot.n}$$

با تعداد $knot.n$ گره:

گام ۱ (شبیه‌سازی هسته و دامنه): هسته‌ی متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی $\tilde{X}_{knot.n}$ از توزیع f_C ، و پهناهای چپ و راست بترتیب از توزیع‌های $f_{S_X^l}$ و $f_{S_X^r}$ شبیه‌سازی شوند. بنابراین، در این گام هسته و دامنه‌ی متغیر فازی $\tilde{X}_{knot.n}$ بترتیب برابر C_x و $[C_x - S_x^l, C_x + S_x^r]$ تعیین/مشاهده/شبیه‌سازی می‌شوند. گام ۲ (شبیه‌سازی گره‌ها): مختصات گره‌های یال چپ، به ازای $i = 1, 2, \dots, knot.n$ ، بصورت $(C_x - l_{(knot.n+1-i)}, \frac{i}{knot.n+1})$ است که در آن $l_{(i)}$ داده‌های ترتیبی شبیه‌سازی شده از توزیع بریده‌شده‌ی $f_{S_X^l}$ در بازه‌ی $[0, S_x^l]$ است، یعنی

$$l_1, \dots, l_{knot.n} \stackrel{iid}{\sim} f_{S_X^l; Trun[0, S_x^l]}$$

مختصات گره‌های یال راست، به ازای $i = 1, 2, \dots, knot.n$ ، بصورت $(C_x + r_{(i)}, \frac{i}{knot.n+1})$ است که در آن $r_{(i)}$ داده‌های ترتیبی شبیه‌سازی شده از توزیع بریده‌شده‌ی $f_{S_X^r}$ در بازه‌ی $[0, S_x^r]$ است $(r_1, \dots, r_{knot.n} \stackrel{iid}{\sim} f_{S_X^r; Trun[0, S_x^r]})$.

گام ۳ (اتصال گره‌ها): پس از اتصال ترتیبی نقاط/گره‌های حاصل در گام‌های ۱ و ۲ بوسیله‌ی چندین پاره‌خط مستقیم، تابع عضویت مشاهده‌ی متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی $\tilde{x}_{knot.n}$ شکل گرفته و شبیه‌سازی متغیر تصادفی فازی تمام می‌شود.

توجه ۲.۲. تعریف متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی، مبتنی بر سه متغیر تصادفی غیرفازی (یکی حقیقی-مقدار و دوتای دیگر نامنفی-مقدار) است و لذا از نام هر سه متغیر تصادفی معمولی (اولی برای هسته‌ی تصادفی، دومی برای پهنای تصادفی سمت چپ و سومی برای پهنای تصادفی سمت راست) در متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی استفاده می‌شود (مثال ۳.۲ را ببینید).

مثال ۳.۲. برای تولید/شبیه‌سازی یک عدد فازی از متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی

$$\tilde{X}_2 \approx [C_X \sim N(5, 2), S_X^l \sim E(3), S_X^r \sim U(0, 1)]_2$$

که آن را اصطلاحاً یک متغیر تصادفی فازی ۲ گره‌ای نرمال-نمایی-یکنواخت به‌ترتیب با پارامترهای ۵، ۲، ۳، ۰ و ۱ می‌نامیم، ابتدا درگام ۱ از متغیرهای تصادفی $C_X \sim N(5, 2)$ ، $S_X^l \sim E(3)$ و $S_X^r \sim U(0, 1)$ به ترتیب هسته و دامنه‌ی $\tilde{x}_2$ را تولید می‌کنیم:

$$\text{core}(\tilde{x}_2) = C_x = 1.717$$

$$\text{supp}(\tilde{x}_2) = [C_x - S_x^l, C_x + S_x^r] = [1.717 - 0.57, 1.717 + 0.186] = [1.147, 1.903]$$

درگام ۲، ابتدا تعداد $knot.n = 2$ داده‌ی تصادفی از توزیع بریده‌شده‌ی $E(3)$ در بازه‌ی $[0, 0.57]$ ، که تابع چگالی احتمال آن بصورت

$$f_{S_X^l; \text{Trun}[0, S_x^l]} = f_{E(3); \text{Trun}[0, 0.57]} = \frac{f_E(3) I(0 < x < 0.57)}{F_E(3)(0.57) - F_E(3)(0)} = \frac{3e^{3x}}{0.95}, \quad 0 < x < 0.57$$

است تولید می‌کنیم ($l_1 = 0.28$ و $l_2 = 0.17$). بنابراین $l(2) = 0.28$ و $l(1) = 0.17$ و مختصات گره‌های شبیه‌سازی‌شده‌ی یال چپ به ازای $i = 1, 2$ برابر

$$\left(C_x - l_{(knot.n+1-i)}, \frac{i}{knot.n+1} \right) = \left(1.717 - l_{(3-i)}, \frac{i}{3} \right)$$

و در نتیجه گره‌های $(1.689, 0.333)$ و $(1.700, 0.667)$ برای یال چپ شبیه‌سازی شدند.

برای تعیین گره‌ها/نقاط روی یال راست نیز ابتدا دو داده‌ی $r_1 = r(1) = 0.52$ و $r_2 = r(2) = 0.156$ از توزیع

$$\begin{aligned} f_{S_X^r; \text{Trun}[0, S_x^r]} &= f_{U(0, 2); \text{Trun}[0, 0.186]} = \frac{f_{U(0, 2)}}{F_{U(0, 2)}(0.186) - F_{U(0, 2)}(0)} I(0 < x < 0.186) \\ &= \frac{1}{2 \times 0.093} I(0 < x < 0.186) = f_{U(0, 0.186)} \end{aligned} \quad (1)$$

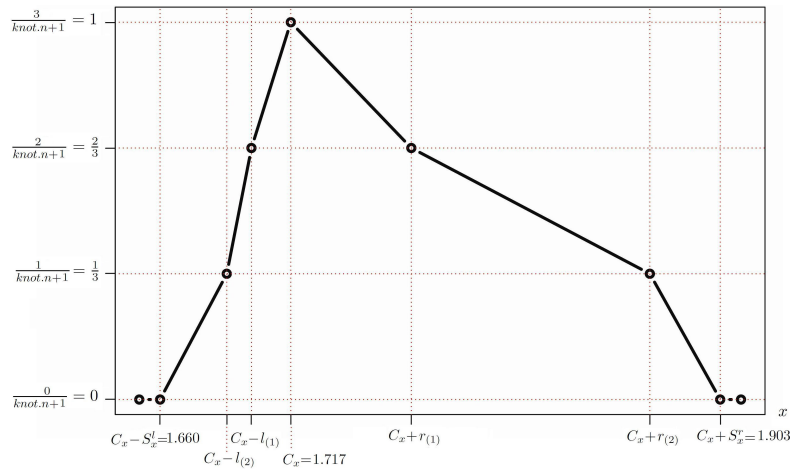
مشاهده شده‌اند. لذا مختصات گره‌های شبیه‌سازی‌شده‌ی یال راست، با توجه به فرمول

$$\left(C_x + r_{(i)}, \frac{i}{knot.n+1} \right) = \left(1.717 + r_{(i)}, \frac{i}{3} \right), \quad i = 1, 2$$

برابر نقاط $(1.873, 0.333)$ و $(1.769, 0.667)$ بدست می‌آیند. سپس درگام ۳، نقاط/گره‌های شبیه‌سازی‌شده درگام‌های ۱ و ۲ را بوسیله‌ی خطوط مستقیم به یکدیگر متصل می‌کنیم تا تابع عضویت $\tilde{x}_2$ همانند شکل ۱ رسم شود.

تعریف ۴.۲. تابع توزیع تجربی پیوسته بر اساس داده‌های نمونه‌ی تصادفی $x_1, \dots, x_n$ از اتصال پی‌درپی نقاط $(x_{(i)}, \frac{i}{n})$ به ازای $i = 0, 1, \dots, n$ بدست می‌آید که در آن $x_{(i)}$ برابر i امین داده‌ی ترتیبی (به ازای $i = 1, \dots, n$) و $x_{(0)}$ کوچک‌ترین نقطه دامنه X است. بنابراین تابع توزیع تجربی پیوسته بر اساس داده‌های $x_1, \dots, x_n$ برابر است با

$$F_{x_1, \dots, x_n}^c(x) = \begin{cases} 0 & \text{اگر } x < x_{(0)} \\ \frac{i}{n} + \frac{x - x_{(i)}}{n(x_{(i+1)} - x_{(i)})} & \text{اگر } x_{(i)} \leq x < x_{(i+1)}, i = 0, 2, \dots, n-1 \\ 1 & \text{اگر } x_{(n)} \leq x \end{cases} \quad (2)$$



شکل ۱: تابع عضویت عدد فازی مشاهده شده از یک متغیر تصادفی قطعه قطعه شدنی خطی با تعداد $knot.n = ۲$ گره در مثال ۳.۲.

نتیجه ۵.۲. با توجه به تعریف ۱.۲ و تعریف ۴.۲، می توان تابع عضویت متغیر تصادفی فازی قطعه قطعه شدنی خطی

$$\tilde{X}_{knot.n} \approx [C_X \sim f_C, S_X^l \sim f_{S_X^l}, S_X^r \sim f_{S_X^r}]_{knot.n}$$

را بر حسب تابع توزیع تجربی پیوسته بصورت زیر بازنویسی کرد

$$\tilde{X}_{knot.n}(x) = \begin{cases} 0 & \text{اگر } x \leq C_x - S_x^l \\ 1 - F_{l_1, \dots, l_{knot.n}}^c(C_x - x) & \text{اگر } C_x - S_x^l < x < C_x \\ 1 & \text{اگر } x = C_x \\ 1 - F_{r_1, \dots, r_{knot.n}}^c(x - C_x) & \text{اگر } C_x < x < C_x + S_x^r \\ 0 & \text{اگر } x \geq C_x + S_x^r \end{cases} \quad (۳)$$

مثال ۶.۲. قصد داریم تا بر اساس نتیجه ۵.۲، تابع عضویت یک مشاهده از متغیر تصادفی فازی قطعه قطعه شدنی خطی

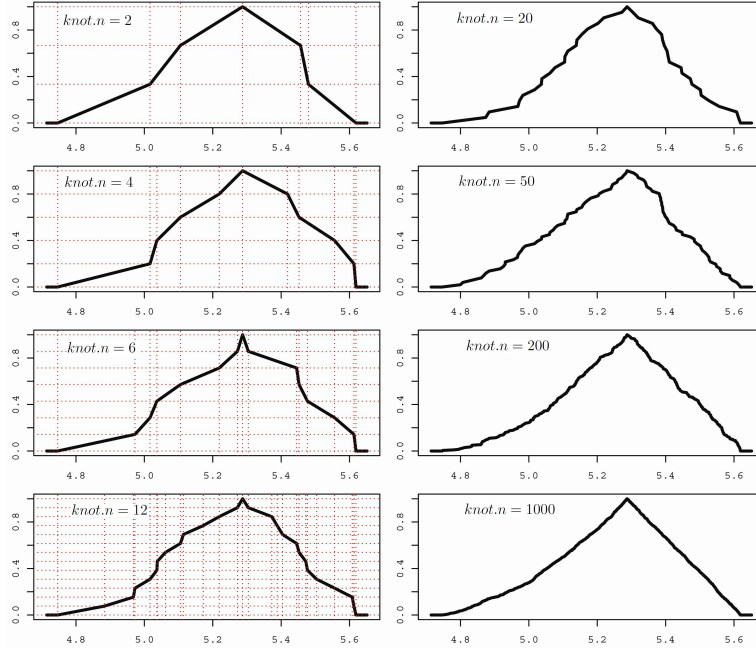
$$\tilde{X}_{knot.n} \approx [C_X \sim N(۳, ۱), S_X^l \sim E(۳), S_X^r \sim U(۰, ۲)]_{knot.n}$$

که اصطلاحاً یک متغیر تصادفی فازی نرمال-نمایی-یکنواخت به ترتیب با پارامترهای ۳، ۱، ۳، ۰ و ۲ با تعداد $knot.n$ گره نامیده می شود را به ازای مقادیر مختلفی از تعداد گره ها شبیه سازی کنیم. تابع عضویت $\tilde{x}_{knot.n}$ به ازای مقادیر ۱۰۰۰، ۲۰۰، ۵۰، ۲۰، ۱۲، ۶، ۴، ۲، $knot.n = ۲$ در شکل ۲ رسم شده است. بطور شهودی، شکل ۲ الهام بخش نوعی همگرایی در حالت حدی $knot.n \rightarrow \infty$ است که در بخش بعدی مورد بررسی قرار گرفته و موجب ارائه تعریفی ساده تر و دقیق تر برای متغیر تصادفی فازی شده است.

۳ متغیر تصادفی فازی LR

تعریف ۱.۳. فرض کنید عدد فازی N یک زیرمجموعه ی فازی از $\mathbb{R}$ با تابع عضویت

$$N(x) = \begin{cases} L\left(\frac{n-x}{\alpha}\right) & \text{اگر } x \leq n \\ R\left(\frac{x-n}{\beta}\right) & \text{اگر } x > n \end{cases} \quad (۴)$$



شکل ۲: تابع عضویت متغیر تصادفی فازی مشاهده شده قطعه قطعه شده خطی در مثال ۶.۲ با افزایش تعداد گره ها.

باشد به طوری که در آن $L: \mathbb{R}^+ \cup \circ \rightarrow [\circ, 1]$ و $R: \mathbb{R}^+ \cup \circ \rightarrow [\circ, 1]$ توابعی غیر نزولی، $R(\circ) = L(\circ) = 1$ و $\alpha, \beta > \circ$ در این صورت $N = (n, \alpha, \beta)_{LR}$ نشان می دهند که در آن L, β, α, n به ترتیب مقدار هسته، پهنای چپ، پهنای راست، تابع شکل چپ و تابع شکل راست نامیده می شوند؛ (۲؛ ۴).

قضیه ۲.۳. در حالت حدی $knot.n \rightarrow \infty$ ، مقدار تابع عضویت متغیر تصادفی قطعه قطعه شده خطی

$$\tilde{X}_{knot.n} \approx [C_X \sim f_C, S_X^l \sim f_{S_X^l}, S_X^r \sim f_{S_X^r}]_{knot.n}$$

(با تعداد گره $knot.n$) در هر نقطه ی حقیقی x ، به

$$\tilde{X}(x) = \begin{cases} \circ & \text{اگر } x \leq C_x - S_x^l \\ 1 - F_{S_X^l; Trun[\circ, S_x^l]}(C_x - x) & \text{اگر } C_x - S_x^l < x < C_x \\ 1 & \text{اگر } x = C_x \\ 1 - F_{S_X^r; Trun[\circ, S_x^r]}(x - C_x) & \text{اگر } C_x < x < C_x + S_x^r \\ \circ & \text{اگر } x \geq C_x + S_x^r \end{cases} \quad (5)$$

با احتمال یک همگراست. به عبارت دیگر به ازای هر $x \in R$

$$\tilde{X}_{knot.n}(x) \xrightarrow{knot.n \rightarrow \infty} \tilde{X}(x). \quad (6)$$

قضیه ۳.۳. در حالت حدی $knot.n \rightarrow \infty$ ، متغیر تصادفی قطعه قطعه شده خطی

$$\tilde{X}_{knot.n} \approx [C_X \sim f_C, S_X^l \sim f_{S_X^l}, S_X^r \sim f_{S_X^r}]_{knot.n}$$

با احتمال یک به متغیر تصادفی فازی $\tilde{X} \approx (C_x, S_x^l, S_x^r)_{LR}$ با توابع شکل

$$L(x) = 1 - F_{S_X^l; Trun[\circ, S_x^l]}(xS_x^l), \quad \circ \leq x \leq 1 \quad (۷)$$

$$R(x) = 1 - F_{S_X^r; Trun[\circ, S_x^r]}(xS_x^r), \quad \circ \leq x \leq 1 \quad (۸)$$

میل می‌کند (که در گام ۱ فیکس/مشاهده شده است)؛ یعنی

$$\tilde{X}_{knot.n} \xrightarrow{knot.n \rightarrow \infty} (C_x, S_x^l, S_x^r)_{LR}. \quad (۹)$$

با توجه به قضیه ۳.۳، می‌توان تعریف ۱.۲ را برای حالت حدی $knot.n \rightarrow \infty$ بصورت زیر بیان کرد. از آنجا که تعداد گره‌ها در تعریف زیر به سمت بینهایت میل می‌کند، لذا بدیهی است که تعریف ۴.۳ تعبیر دقیق‌تری از یک متغیر تصادفی فازی (نسبت به تعریف ۱.۲) دارد که در ادامه آن را متغیر تصادفی فازی LR می‌نامیم.

تعریف ۴.۳. $\tilde{X} \approx (C_X \sim f_C, S_X^l \sim f_{S_X^l}, S_X^r \sim f_{S_X^r})_{LR}$ را یک متغیر تصادفی فازی-مقدار LR (متغیر تصادفی فازی LR) نامیم، هرگاه: (الف) هسته‌ی آن یک متغیر تصادفی حقیقی-مقدار از توزیع $C_X \sim f_C$ و پهنای‌های چپ و راست آن متغیرهای تصادفی نامنفی-مقدار از توزیع‌های $S_X^l \sim f_{S_X^l}$ و $S_X^r \sim f_{S_X^r}$ باشند؛ (ب) توابع شکل بر حسب توابع توزیع بریده شده، بصورت زیر باشند

$$\begin{aligned} L(x) &= 1 - F_{S_X^l; Trun[\circ, S_x^l]}(xS_x^l), \quad \circ \leq x \leq 1, \\ R(x) &= 1 - F_{S_X^r; Trun[\circ, S_x^r]}(xS_x^r), \quad \circ \leq x \leq 1. \end{aligned} \quad (۱۰)$$

تعریف ۵.۳. $\tilde{X} \approx (C_X \sim f_C, S_X \sim f_{S_X})_L$ را یک متغیر تصادفی فازی متقارن L نامیم، هرگاه: (الف) هسته‌ی آن یک متغیر تصادفی حقیقی-مقدار از توزیع $C_X \sim f_C$ و پهنای آن متغیرهای تصادفی نامنفی-مقدار از توزیع $S_X \sim f_{S_X}$ باشند؛ (ب) تابع شکل بر حسب تابع توزیع بریده شده، بصورت زیر باشد

$$L(x) = 1 - F_{S_X; Trun[\circ, S_x]}(xS_x), \quad \circ \leq x \leq 1. \quad (۱۱)$$

۴ متغیر تصادفی فازی مثلثی

قضیه ۱.۴. در حالت حدی $knot.n \rightarrow \infty$ ، متغیر تصادفی قطعه‌قطعه شده‌ی خطی

$$\tilde{X}_{knot.n} \approx [C_X \sim f_C, S_X^l \sim U(\circ, l), S_X^r \sim U(\circ, r)]_{knot.n}, \quad l, r > \circ$$

با احتمال یک به متغیر تصادفی فازی مثلثی $\tilde{X} \approx (C_x, S_x^l, S_x^r)_T$ که در گام ۱ فیکس/مشاهده شده است میل می‌کند؛ یعنی

$$[C_X \sim f_C, S_X^l \sim U(\circ, l), S_X^r \sim U(\circ, r)]_{knot.n} \xrightarrow{knot.n \rightarrow \infty} (C_x, S_x^l, S_x^r)_T. \quad (۱۲)$$

نتیجه ۲.۴. برای شبیه‌سازی متغیر تصادفی فازی مثلی $\tilde{X}_T \approx (C_x, S_x^l, S_x^r)_T$ ، در حالت حدی $knot.n \rightarrow \infty$ می‌توان در تعریف ۱.۲ از گام ۲ صرف‌نظر کرد. به عبارت ساده‌تر، برای تولید/شبیه‌سازی متغیر تصادفی فازی مثلی $\tilde{X} \approx (C_x, S_x^l, S_x^r)_T$ تنها کافیت هسته‌ی متغیر تصادفی فازی مثلی را از توزیع f_C و پهناهای چپ و راست آن را بترتیب از توزیع‌های $f_{S_X^l}$ و $f_{S_X^r}$ شبیه‌سازی کرد. بنابراین، تابع عضویت متغیر تصادفی فازی مثلی $\tilde{X}_T$ بصورت

$$\tilde{X}_T(x) = \begin{cases} \frac{x-C_X+S_X^l}{S_X^l} & \text{اگر } S_X^l - C_X \leq x < C_X \\ \frac{C_X+S_X^r-x}{S_X^r} & \text{اگر } C_X \leq x < C_X + S_X^r \\ 0 & \text{جاهای دیگر} \end{cases} \quad (۱۳)$$

است که در آن $C_X \sim f_C$ ، $S_X^l \sim f_{S_X^l}$ و $S_X^r \sim f_{S_X^r}$.

نتیجه‌گیری و بحث

در این مقاله، ابتدا یک تعریف منطقی برای متغیر تصادفی فازی، تحت عنوان «متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی» پیشنهاد شده است (تعریف ۱.۲). سپس با بررسی حالت حدی مطرح شده در دومین گام تولید متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی، تعریفی ساده‌تر و در عین حال دقیق‌تر (نسبت به تعریف متغیر تصادفی فازی قطعه‌قطعه‌شده‌ی خطی) مطرح شده که در ادامه آن را، بعنوان اصلی‌ترین نتیجه‌ی این مقاله، به اختصار بازگو می‌سازیم. $(C_X \sim f_C, S_X^l \sim f_{S_X^l}, S_X^r \sim f_{S_X^r})_{LR}$ را یک «متغیر تصادفی فازی LR» گوئیم، هرگاه هسته‌ی آن یک متغیر تصادفی حقیقی-مقدار از توزیع f_C ، $C_X \sim f_C$ ، پهناهای چپ و راست آن متغیرهای تصادفی نامنفی-مقدار از توزیع‌های $f_{S_X^l}$ و $f_{S_X^r}$ ، و همچنین توابع شکل به ازای $0 \leq x \leq 1$ بصورت $L(x) = F_{S_X^l; Trun[0, S_x^l]}(S_x^l(1-x))$ و $R(x) = 1 - F_{S_X^r; Trun[0, S_x^r]}(xS_x^r)$ در نظر گرفته شوند.

بنابراین می‌توان ادعا کرد که متغیر تصادفی فازی را فقط مقادیر مشاهده‌شده‌ی C_x ، S_x^l و S_x^r (در گام ۱) و همچنین توابع توزیعی که یال‌های چپ و راست را می‌سازند، رقم می‌زنند و نه مقدار $knot.n$ (تعداد گره‌های تصادفی) آن‌ها. به عبارت دیگر، تصادفی بودن یک متغیر تصادفی فازی LR تنها ریشه در تصادفی بودن هسته و دامنه/سپوریت آن دارد.

تعریف متغیر تصادفی فازی LR در بر دارنده‌ی دو نوع عدم قطعیت احتمالی و امکانی بطور همزمان است. ویژگی متغیرهای تصادفی فازی LR در نوع تابع عضویت آنهاست و محاسبه‌ی توابع/آماره‌ها، که از اعمال عملگرهای حسابی بر روی این متغیرها بوجود می‌آیند، از قواعد خاصی پیروی می‌کند. این ویژگی باعث می‌شود که در بسیاری از مباحث آمار فازی (همچون آزمون فرضیه‌ها در محیط فازی، رگرسیون فازی، فاصله اطمینان فازی، کنترل کیفیت فازی، قضایای حدی برای متغیرهای تصادفی فازی و ...) از این نوع متغیرهای تصادفی فازی استفاده شود.

مراجع

- [۱] چاچی، ج.، روش‌های آماری بر اساس اطلاعات نادقیق، رساله دکتری آمار، دانشکده علوم ریاضی، دانشگاه صنعتی اصفهان، ۱۳۹۱.
- [۲] طاهری، س.م.، ماشین‌چی، م.، مقدمه‌ای بر آمار و احتمال فازی، انتشارات دانشگاه شهید باهنر کرمان، چاپ اول، ۱۳۸۷.

[4] Dubois, D. and Prade, H. (1980), *Fuzzy Sets and Systems: Theory and Applications*. Academic Press.

[4] Nourbakhsh, M.R., Mashinchi, M. and Parchami, A. (2013), Analysis of variance based on fuzzy observations. *International Journal of Systems Science* 44, 714-726.

- [5] Puri, M. L. and Ralescu, D. A. (1985) The concept of normality for fuzzy random variables. *The Annals of Probability* 13, 1373-1379.
- [6] Puri, M. L. and Ralescu, D. A. (1986) Fuzzy random variables. *Journal of Mathematical Analysis and Applications* 114, 409-422.
- [7] Wu, H.C. (1999), The central limit theorems for fuzzy random variables. *Information Sciences* 120, 239-256.
- [8] Wu, H.C. (2000), The laws of large numbers for fuzzy random variables. *Fuzzy Sets and Systems* 116, 245-262.

رویکردی جدید به منظور مدل‌سازی و حل مسأله معکوس ۱- میانه فازی

منا خداقلی، دکتر اردشیر دولتی، علی حسین زاده^۱

گروه ریاضی دانشگاه شاهد

چکیده

مسائل مکان‌یابی تأسیسات و تسهیلات، از کاربردی‌ترین مسائل شاخه تحقیق در عملیات به‌شمار می‌رود. از کاربردهای معروف این مسأله می‌توان به مکان‌یابی انبارها، بیمارستان‌ها، ایستگاه‌های امداد و نجات، بانک‌ها، تأسیسات نظامی و ... اشاره کرد. هدف از حل این مسائل تعیین مکان مناسبی جهت استقرار تسهیلات و مراکز خدمات رسانی به مشتریان است، به گونه‌ای که این مراکز حداکثر کارایی و خدمات رسانی را به سایر مشتریان داشته باشند. اخیراً مسائل مکان‌یابی با رویکرد معکوس نیز مورد مطالعه و بررسی بسیاری از محققین قرار گرفته است. یکی از معروف‌ترین و کاربردی‌ترین توابع هدف مکان‌یابی که پارامتر تجمع را نیز در نظر می‌گیرد، تابع هدف ۱- میانه می‌باشد. همچنین الگوریتم‌های کلاسیک متنوعی برای معکوس این مسأله معرفی شده است. اما از آن جا که در دنیای واقعی پارامترهای مسأله نادقیق و غیرقطعی هستند، در این مقاله معکوس مسأله ۱- میانه را در حالت فازی روی درخت‌ها مورد مطالعه قرار می‌دهیم و با استفاده از روش‌های حل برنامه‌ریزی فازی به حل آن می‌پردازیم. در انتها نیز مثالی عددی برای اثبات الگوریتم حل ارائه می‌نماییم.

کلمات کلیدی: مکان‌یابی، معکوس مسأله ۱- میانه، درخت، شرط بهینگی، برنامه‌ریزی خطی فازی

۱ مقدمه

مسائل مکان‌یابی و تحلیل موقعیت تأسیسات و تسهیلات، از مهم‌ترین مسائل بهینه‌سازی است که اخیراً معکوس آنها نیز مورد توجه محققین فراوانی قرار گرفته است. این مسائل دارای کاربردهای زیادی در زمینه استقرار تسهیلات در بخش درمانی، اداری، نظامی، آموزشی، برنامه‌ریزی شهری و استقرار ایستگاه‌های آتش‌نشانی، استقرار نیروگاه‌های برق و ... است. هدف از حل این نوع مسائل، تعیین مکانی مناسب جهت احداث مراکز خدمات‌رسانی و استقرار تسهیلات می‌باشد، به گونه‌ای که این مراکز حداکثر کارایی و خدمات رسانی را به سایر مشتریان داشته باشند. در مسائل مکان‌یابی کلاسیک، مکان‌های بهینه جهت استقرار تسهیلات جدید با توجه به توابع

^۱ نام ارائه دهنده مقاله : alinevmath@chmail.ir

هدف، مشخص می‌شوند. دو تابع هدف معروف که به طور معمول ممکن است مورد استفاده قرار گیرد تابع میانه و تابع مرکز می‌باشد. میانه، رأس یا رئوسی از گراف است که مجموع همه فاصله‌های وزنی آنها تا سایر رئوس مینیمم باشد. مرکز نیز رأس یا رئوسی از گراف است که طولانی‌ترین فاصله از آنها تا سایر گره‌ها مینیمم است. در این مقاله، معکوس مسأله ۱- میانه روی درخت را مورد بررسی قرار خواهیم داد. هدف از حل این مسأله، تغییر برخی پارامترهای مسأله مانند وزن رئوس است، به گونه‌ای که رأس از پیش مشخص شده‌ای با توجه به پارامترهای تغییر یافته، به یک رأس ۱-میانه تبدیل شود. این مسأله توسط الگوریتمی که (۵) ارائه داد، در زمان چند جمله‌ای $O(n \log n)$ قابل حل است. با توجه به مطالعات ما در این راستا، از آن جا که معکوس مسأله ۱-میانه با پارامترهای غیر قطعی تحت نظریه فازی تا به حال مورد بررسی قرار نگرفته است، انگیزه‌ای شد که در این مقاله به بررسی این مسأله با هزینه اصلاحات و وزن رئوس غیر قطعی بپردازیم. اگر وزن هر گره، میزان تقاضای مشتریان را نشان دهد، هدف از حل معکوس مسأله ۱- میانه روی درخت‌ها، اصلاح میزان تقاضای مشتریان با حداقل هزینه تحت اصلاحات کراندار است، به گونه‌ای که مکان (رأس) از پیش مشخص شده‌ای، مکان بهینه (۱-میانه) شود. اما از آنجا که در این مسأله میزان تقاضای مشتریان به مراکز خدمات رسانی و تسهیلات با عدم قطعیت روبه رو می‌باشد، بنابراین بدست آوردن میزان دقیق تقاضای بهینه مشتریان به مراکز تسهیلات به دور از واقعیت خواهد بود. بدین منظور ما در این مقاله به منظور رفع عدم قطعیت، از روشهای فازی برای حل مسأله استفاده می‌کنیم. مجموعه‌های فازی در ابتدا توسط (۱۵) معرفی شد. (۳) تصمیم‌گیری در محیط فازی را بیان کردند. ایده بلمن و زاده در تصمیم‌گیری برای اولین بار توسط (۱۴)، در حوزه برنامه‌ریزی ریاضی به کار گرفته شد. یکی از مهم‌ترین کاربردهای تصمیم‌گیری فازی، برنامه‌ریزی خطی فازی می‌باشد که کارهای زیادی در این زمینه انجام شده است. (۴) روشی را برای پیدا کردن جواب فازی مسأله برنامه‌ریزی خطی با تغییر تابع هدف به یک مسأله چند هدفه را معرفی کردند. (۱۳) روش حلی را برای مسائل LP با ضرایب هدف فازی پیشنهاد دادند. (۱) نیز روشی را بر پایه غیرفازی سازی برای حل مسأله برنامه‌ریزی خطی تماماً فازی (FFLP) پیشنهاد دادند. (۱۰) با استفاده از تابع RANK و اعداد فازی L-R روشی نوین برای حل FFLP ارائه کردند. در زمینه مدل‌سازی مکان‌یابی با داده‌های غیرقطعی نیز با استفاده از منطق فازی کارهای زیادی انجام شده است. (۱۷) مسائل مکان‌یابی با وزن طول‌های فازی را در نظر گرفته‌اند و روشی برای حل آن ارائه داده‌اند. (۱۲) نیز یک فرمول‌بندی برای یافتن جواب بهینه مسأله p - میانه در محیط فازی ارائه داده‌اند. (۱۳) نیز به حل مسأله p - میانه با تبدیل به حالت Crisp پرداخته‌اند. بنابراین سازماندهی مقاله بصورت زیر می‌باشد:

در بخش دوم به ارائه برخی از تعاریف و مفاهیم اولیه نظریه مجموعه فازی می‌پردازیم. در بخش سوم معکوس مسأله ۱-میانه را بیان می‌کنیم و در بخش چهارم این مسأله را در حالت فازی بیان می‌کنیم و برای حل آنها از روش‌های برنامه‌ریزی خطی فازی استفاده می‌کنیم. در آخر نیز جهت فهم بهتر مدل، یک مثال عددی ارائه می‌کنیم.

۲ تعاریف اولیه

در این بخش برخی مفاهیم پایه‌ای و خواص نظریه مجموعه فازی مورد نیاز مقاله بیان می‌شود.

تعریف ۱.۰۲. اعداد فازی^۱ و ویژگی‌های آن. عدد فازی A ، یک مجموعه فازی روی خط اعداد حقیقی است که در دو شرط نرمال بودن و تحدب صدق نماید. از آنجا که اکثر مجموعه‌های فازی شرایط نرمال بودن و تحدب را برآورده می‌سازند، بنابراین اعداد فازی مرسوم‌ترین نوع مجموعه‌های فازی می‌باشند. اعداد فازی دارای انواع مختلفی می‌باشند. یکی از آنها عدد فازی مثلثی می‌باشد که در ادامه ویژگی‌های آن مورد بحث قرار می‌گیرند.

تعریف ۲.۰۲. عدد فازی مثلثی. یک عدد فازی $\tilde{A} = (a, b, c)$ یک عدد فازی مثلثی است که a شاخه چپ، b میانه و c شاخه راست آن می‌باشد. یک عدد فازی مثلثی مثبت است اگر فقط اگر $a \geq 0$.

گزاره ۳.۰۲. دو عدد فازی مثلثی $\tilde{A} = (a, b, c)$ و $\tilde{B} = (e, f, g)$ معادلند اگر فقط اگر $a = e$ و $b = f$ و $c = g$ باشد.

تعریف ۴.۰۲. تابع رتبه. فرض کنید $F(\mathbb{R})$ یک مجموعه از اعداد فازی تعریف شده روی مجموعه اعداد حقیقی باشد. در این صورت $R : F(\mathbb{R}) \rightarrow \mathbb{R}$ یک تابع رتبه است که هر عدد فازی را به محور حقیقی (که ترتیب در آن موجود است) نگاشت می‌کند.

اگر $\tilde{A} = (a, b, c)$ یک عدد فازی مثلثی باشد، در این صورت تابع رتبه این عدد فازی به صورت $R(\tilde{A}) = \frac{a + 2b + c}{4}$ می‌باشد.

^۱Fuzzy Numbers

تعریف ۵.۲. عملگرهای حسابی/اعداد مثلثی. در این قسمت عملگرهای حسابی بین دو عدد فازی مثلثی که روی یک مجموعه کلی از اعداد حقیقی $\mathbb{R}$ می باشد بررسی شده است. فرض کنید $\tilde{A} = (a, b, c)$ و $\tilde{B} = (e, f, g)$ دو عدد فازی مثلثی باشند، در این صورت داریم:

$$۱. \tilde{A} \oplus \tilde{B} = (a, b, c) \oplus (e, f, g) = (a + e, b + f, c + g)$$

$$۲. -\tilde{A} = -(a, b, c) = (-a, -b, -c)$$

$$۳. \tilde{A} \ominus \tilde{B} = (a, b, c) \ominus (e, f, g) = (a - e, b - f, c - g)$$

$$۴. \tilde{A} \otimes \tilde{B} = \begin{cases} (ae, bf, cg) & a \geq 0 \\ (ag, bf, cg) & a < 0, c \geq 0 \\ (ag, bf, ce) & c < 0 \end{cases}$$

تعریف ۶.۲. رتبه بندی لکسیکوگرافی. فرض کنید $\tilde{u} = (x_1, y_1, z_1)$ و $\tilde{v} = (x_2, y_2, z_2)$ دو عدد فازی مثلثی دلخواه باشند. در این صورت $\tilde{u}$ نسبتاً کوچکتر از $\tilde{v}$ و با $\tilde{u} < \tilde{v}$ نمایش می دهیم اگر فقط اگر:

$$۱. y_1 < y_2$$

$$۲. y_1 = y_2 \text{ و } (z_1 - x_1) > (z_2 - x_2)$$

$$۳. y_1 = y_2 \text{ و } (z_1 - x_1) = (z_2 - x_2) \text{ و } (z_1 + x_1) < (z_2 + x_2)$$

توجه کنید واضح است که $y_1 = y_2$ و $(z_1 - x_1) = (z_2 - x_2)$ و $(z_1 + x_1) = (z_2 + x_2)$ ، اگر فقط اگر $\tilde{u} = \tilde{v}$.

۳ بیان معکوس مسأله ۱- میانه

یک شبکه درخت به صورت $T = (V, E)$ را در نظر بگیرید، به طوریکه $V = \{v_1, \dots, v_n\}$ مجموعه متناهی از رئوس و $E = \{e_1, \dots, e_m\}$ مجموعه یالهای درخت T است. هر یال $e_i \in E$ دارای طول مثبت l_i و هر رأس $v_j \in V$ دارای وزن w_j که $w_j \in \mathbb{R}^+$ است. طول کوتاهترین مسیر بین دو گره $u, v \in V$ با $d(u, v)$ نمایش داده می شود. فرض کنید $W(T)$ مجموع وزن همه رئوس درخت باشد، یعنی $W(T) := \sum_{j=1}^n w_j$. هدف از حل معکوس مسأله ۱- میانه با تغییر وزن رئوس روی درختها، اصلاح وزن رئوس با کمترین هزینه به صورت w^* ، تحت تغییرات کراندار است، به گونه ای که رأس از پیش مشخص شده v_s که نسبت به مجموعه وزن رئوس داده شده w ، ۱- میانه نیست، به یک رأس ۱- میانه تبدیل شود. اگر افزایش و کاهش وزن هر رأس $v_j \in V$ را به ترتیب با p_j و q_j نشان دهیم، آنگاه وزن هر رأس از درخت داده شده به صورت $w_j^* = w_j + p_j - q_j$ (برای $j = 1, \dots, n$) اصلاح می شود. (۱۱) و یک سال بعد (۹) نشان دادند که جواب مسأله ۱- میانه روی درختها کاملاً از طولهای یالی مستقل بوده و فقط به وزنهای رأسی وابسته است.

لم ۱۰.۳. (۹) و (۱۱) (شرط بهینگی (۱- میانه بودن یک رأس)): درخت $T = (V, E)$ را در نظر بگیرید. رأس $v_s \in V$ از درجه k ، یک ۱- میانه درخت T است، اگر و فقط اگر $W(T_i) \leq \frac{W(T)}{۲}$ ، برای $i = 1, \dots, k$ به طوری که $W(T_i) := \sum_{v_j \in T_i} w_j$ وزن هر زیردرخت ماکسیمال حاصل از v_s است.

لم ۲۰.۳. (۸) فرض کنید $W(T_1) \leq \frac{W(T)}{۲}, \dots, W(T_{k-1}) \leq \frac{W(T)}{۲}$ و $W(T_k) > \frac{W(T)}{۲}$ در این صورت، رأس $v_s \in V$ یک ۱- میانه درخت T نسبت به وزنهای اصلاح شده است اگر و فقط اگر:

$$D := W(T_k) - \frac{W(T)}{۲} = 0$$

بنابراین مدل ریاضی این مسأله که توسط (۵) و (۸) ارائه شده است، به صورت زیر است:

$$\min \sum_{j=1}^n c_j x_j \quad (۱)$$

$$s.t \quad \sum_{j=1}^n x_j = b \quad (2)$$

$$\circ \leq x_j \leq \bar{x}_j \quad \forall \quad j = 1, \dots, n \quad (3)$$

$$\bar{x}_j := \begin{cases} \bar{p}_j & \text{اگر } x_j = p_j \\ \bar{q}_j & \text{اگر } x_j = q_j \end{cases} \text{ و } x_j := \begin{cases} q_j & v_j \in T_k \\ p_j & v_j \notin T_k \end{cases}$$

که $b := 2D = 2W(T_k) - W(T)$ و c_j مثبت هستند و به طوریکه $\bar{x}_j$ کران بالای تغییرات وزن رئوس را نشان می دهد.

مسئله بالا یک مسئله کوله پشتی پیوسته است که توسط الگوریتم (۲) در زمان خطی قابل حل است. بنابراین معکوس مسئله ۱- میانه روی درخت با امکان تغییر وزن رئوس (با وزن های مثبت)، در زمان $O(n)$ قابل حل است.

۴ مدل سازی معکوس مسئله ۱- میانه فازی و الگوریتم حل آن

یک مسئله برنامه ریزی خطی با m محدودیت فازی و n متغیر فازی و با ضرایب فازی تابع هدف به صورت زیر فرمول بندی می شود:

$$Max(Min) \tilde{C}^T \otimes \tilde{X} \quad (4)$$

$$\tilde{A} \otimes \tilde{X} = \tilde{B} \quad (5)$$

$$\tilde{X} \in F^+ \quad (6)$$

به طوریکه $\tilde{C}^T = [\tilde{c}_j]_{1 \times n}$, $\tilde{X} = [\tilde{x}_j]_{n \times 1}$, $\tilde{A} = [\tilde{a}_{ij}]_{m \times n}$, $\tilde{B} = [\tilde{b}_i]_{m \times 1}$ با فرض اینکه هزینه اصلاح وزن رئوس و هم چنین وزن رئوس فازی هستند، ابتدا معکوس مسئله ۱- میانه را به مدل فازی تبدیل کرده و سپس آن را با استفاده از تابع Rank به حالت CRISP تبدیل کرده و با استفاده از روشهای معمول به حل آن می پردازیم. در این صورت با قرار دادن $\tilde{C}^T = [\tilde{c}_j]_{1 \times n}$, $\tilde{X} = [\tilde{x}_j]_{n \times 1}$, $\tilde{B} = [\tilde{b}]_{1 \times 1}$ و $\tilde{A} = \tilde{I} = [\tilde{1}]_{1 \times n}$ در مدل ۴-۶، مسئله کلاسیک معکوس ۱- میانه به مسئله برنامه ریزی خطی تماماً فازی به صورت زیر فرموله می شود:

$$\min \sum_{j=1}^n \tilde{c}_j \tilde{x}_j \quad (7)$$

$$s.t \quad \sum_{j=1}^n \tilde{x}_j = \tilde{b} \quad (8)$$

$$\circ \leq \tilde{x}_j \leq \tilde{\bar{x}}_j \quad \forall \quad j = 1, \dots, n \quad (9)$$

$$\tilde{x}_j, \tilde{b}, \tilde{c}_j, \tilde{\bar{x}}_j \in Triangle \text{ fuzzy number}(TF^+) \quad (10)$$

اکنون برای حل مدل فازی ۷-۱۰ الگوریتم زیر را ارائه می کنیم:

الگوریتم ۱. این الگوریتم دارای گام های زیر است:

- گام ۱. مدل برنامه ریزی خطی تماماً فازی ۷-۱۰ را در نظر بگیرید. به طوریکه اعداد فازی در نظر گرفته شده در مدل اعداد فازی مثلثی باشند.
- گام ۲. با توجه به اینکه در گام ۱ در نظر گرفتیم که $\tilde{x}_j, \tilde{\bar{x}}_j, \tilde{c}_j, \tilde{b} \in TF^+$ ، در این صورت پارامترهای فازی $\tilde{b}, \tilde{c}_j, \tilde{x}_j$ و $\tilde{\bar{x}}_j$ را به ترتیب با اعداد فازی مثلثی (f, g, h) ، (a_j, b_j, c_j) ، (x_j, y_j, z_j) و $(\bar{x}_j, \bar{y}_j, \bar{z}_j)$ نشان می دهیم. سپس مسئله $FFLP$ اشاره شده در گام ۱ به صورت زیر فرموله می شود:

$$\min \sum_{j=1}^n (a_j, b_j, c_j) \otimes (x_j, y_j, z_j) \quad (11)$$

$$s.t \quad \sum_{j=1}^n (x_j, y_j, z_j) = (f, g, h) \quad (12)$$

$$(x_j, y_j, z_j) \leq (\bar{x}_j, \bar{y}_j, \bar{z}_j) \quad \forall \quad j = 1, \dots, n \quad (13)$$

$$(x_j, y_j, z_j) \in TF^+ \quad \forall \quad j = 1, \dots, n \quad (14)$$

• گام ۳. با استفاده از عملگرهای حسابی تعریف شده بر روی اعداد فازی مثالی مسئله $FFLP$ اشاره شده در گام ۲ به شکل مسئله برنامه ریزی خطی کلاسیک (CLP) زیر تبدیل می شود:

$$\min \mathcal{R} \left(\sum_{j=1}^n (a_j, b_j, c_j) \otimes (x_j, y_j, z_j) \right) \quad (15)$$

$$s.t \quad \sum_{j=1}^n x_j = f \quad (16)$$

$$\sum_{j=1}^n y_j = g \quad (17)$$

$$\sum_{j=1}^n z_j = h \quad (18)$$

$$(x_j, y_j, z_j) \leq (\bar{x}_j, \bar{y}_j, \bar{z}_j) \quad \forall \quad j = 1, \dots, n \quad (19)$$

$$y_j - x_j \geq 0, z_j - y_j \geq 0 \quad \forall \quad j = 1, \dots, n. \quad (20)$$

• گام ۴. جواب بهینه x_j^*, y_j^*, z_j^* با استفاده از حل مسئله برنامه ریزی خطی کلاسیک اشاره شده در گام ۳ بدست می آید.

• گام ۵. جواب بهینه فازی با قرار دادن x_j^*, y_j^*, z_j^* در $\tilde{x}_j = (x_j, y_j, z_j)$ بدست می آید.

• گام ۶. مقدار بهینه فازی با قرار دادن در $\sum_{j=1}^n \tilde{c}_j \tilde{x}_j^*$ بدست می آید.

۵ مثال عددی

برای فهم بهتر الگوریتم به حل مثال عددی زیر می پردازیم. در این مدل فرض کرده ایم که بردار ضرایب هزینه ها مقادیری قطعی باشند.

$$\text{Min} Z = 6/325 \times \tilde{x}_1 + 5/656 \times \tilde{x}_2 + 10/816 \times \tilde{x}_3$$

$$S.T \quad (\tilde{x}_{11}, \tilde{x}_{12}, \tilde{x}_{13}) \oplus (\tilde{x}_{21}, \tilde{x}_{22}, \tilde{x}_{23}) \oplus (\tilde{x}_{31}, \tilde{x}_{32}, \tilde{x}_{33}) = (14, 15, 15/5)$$

$$(\tilde{x}_{11}, \tilde{x}_{12}, \tilde{x}_{13}) \leq (4, 5, 6)$$

$$(\tilde{x}_{21}, \tilde{x}_{22}, \tilde{x}_{23}) \leq (6/5, 8, 10)$$

$$(\tilde{x}_{31}, \tilde{x}_{32}, \tilde{x}_{33}) \leq (11/5, 12, 13)$$

$$\tilde{x}_1, \tilde{x}_2, \tilde{x}_3 \in TF^+$$

با توجه به روند الگوریتم پیشنهادی جواب بهینه مسئله بالا برابر است با $\tilde{x}_1 = (4, 5, 5)$, $\tilde{x}_2 = (8, 8, 8)$ و $\tilde{x}_3 = (2, 2, 2)$ و مقدار جواب بهینه تابع هدف برابر $97/63075$ می باشد.

نتیجه‌گیری

مسائل مکان‌یابی فازی یکی از مسائل جدید می‌باشد که می‌تواند به عنوان زمینه‌ای مهم در تحقیقات قرار گیرد. در این مقاله ابتدا یکی از مسائل مهم در مکان‌یابی یعنی مسئله معکوس ۱-میانه را بیان کردیم و با توجه به اینکه داده‌ها در محیط دنیای واقعی دارای ابهام بوده و معمولاً با عدم قطعیت روبرو می‌شوند مسئله را با داده‌ها و متغیرهای فازی مدلسازی کرده و سپس مدل بدست آمده را با تکنیک‌هایی مبتنی بر برنامه ریزی خطی فازی به مساله Crisp تبدیل کرده و آن را حل می‌نماییم.

مراجع

- [1] Allahviranloo, T., Lotfi, F. H., Kiasary, M. K., Kiani, N. A., & Alizadeh, L. (2008). Solving fully fuzzy linear programming problem by the ranking function. *Applied mathematical sciences*, 2(1), 19-32.
- [2] Balas, E., & Zemel, E. (1980). An algorithm for large zero-one knapsack problems. *operations esearch*, 28(5), 1130-1154.
- [3] Bellman, R. E., Zadeh, L. A. (1970). Decision making in a fuzzy environment, *Management Science* 17, 141-164.
- [4] Buckley, J. J., & Feuring, T. (2000). Evolutionary algorithm solution to fuzzy problems: fuzzy linear programming. *Fuzzy sets and systems*, 109(1), 35-53.
- [5] Burkard, R. E., Pleschiutchnig, C., & Zhang, J. (2004). Inverse median problems. *Discrete Optimization*, 1(1), 23-39.
- [6] Dehghan, M., Hashemi, B., & Ghatee, M. (2006). Computational methods for solving fully fuzzy linear systems. *Applied mathematics and Computation*, 179(1), 328-343.
- [7] Ezzati, R., Khorram, E., & Enayati, R. (2015). A new algorithm to solve fully fuzzy linear programming problems using the MOLP problem. *Applied mathematical modelling*, 39(12), 3183-3193.
- [8] Galavii, M. (2010). The inverse 1-median problem on a tree and on a path. *Electronic Notes in Discrete Mathematics*, 36, 1241-1248.
- [9] Goldman, A. J. (1971). Optimal center location in simple networks. *Transportation science*, 5(2), 212-221.
- [10] Hosseinzadeh, A., & Edalatpanah, S. A. (2016). A new approach for solving fully fuzzy linear programming by using the lexicography method. *Advances in Fuzzy Systems*, 2016, 4.
- [11] Hua, L. K. (1962). Application of mathematical models to wheat harvesting. *Chinese Mathematics*, 2, 539-560.
- [12] Kutangila-Mayoya, D., & Verdegay, J. L. (2005). p-Median problems in a fuzzy environment. *Mathware & soft computing*, 2005, vol. 12, núm. 2.
- [13] Taleshian, F., & Fathali, J. (2016). A Mathematical Model for Fuzzy-Median Problem with Fuzzy Weights and Variables. *Advances in Operations Research*, 2016.

- [14] Tanaka, H., Okuda, T., & Asai, K. (1973). Fuzzy mathematical programming. Transactions of the Society of Instrument and Control Engineers, 9(5), 607-613.
- [15] Zadeh, L. A. (1965). Fuzzy sets, Information and Control, 8, 69-78.
- [13] Zhang, G., Wu, Y. H., Remias, M., & Lu, J. (2003). Formulation of fuzzy linear programming problems as four-objective constrained optimization problems. Applied Mathematics and Computation, 139(2-3), 383-399.
- [17] Perez, J. A. M., Vega, J. M. M., & Verdegay, J. L. (2004). Fuzzy location problems on networks. Fuzzy Sets and systems, 142(3), 393-405.

رگرسیون خطی فازی استوار براساس فاصله علامتدار بین دو عدد فازی

امیر حمزه خمر،^۱ محسن عارفی^۲

دانشکده ریاضی و آمار دانشگاه بیرجند، بیرجند، ایران

چکیده

در این مقاله یک روش استوارسازی برای مدل‌های رگرسیون خطی فازی براساس M -برآوردها و فاصله علامتدار بین دو عدد فازی معرفی می‌شود که در حالت خاص یک مدل رگرسیون چندکی فازی استوار ارائه می‌کند. برای معرفی روش استوارسازی، یک مدل رگرسیونی خطی فازی شامل داده‌های ورودی و خروجی فازی مثلثی متقارن و ضرایب رگرسیونی غیرفازی در نظر می‌گیریم. مثال عددی نشان می‌دهد که روش پیشنهادی برخلاف روش حداقل مربعات نسبت به داده‌های پرت حساس نمی‌باشد. در پایان، با استفاده از شاخص نیکویی برازش یک مطالعه مقایسه‌ای با سایر روش‌ها ارائه می‌شود.

کلمات کلیدی: فاصله علامتدار، رگرسیون خطی فازی استوار، عدد فازی مثلثی متقارن، داده پرت، M -برآوردها

۱ مقدمه

تحلیل رگرسیونی یکی از ابزارهای آماری برای تحلیل داده و یافتن روابط بین متغیرهای مستقل و متغیر وابسته می‌باشد. روشهای کلاسیک مبتنی بر مفروضاتی از قبیل دقیق بودن مشاهدات، نرمال بودن و ناهمبسته بودن جملات خطا و ... استوار است. اما در واقعیت با برخی از محدودیت‌ها و مشاهدات مواجه هستیم که این مفروضات برقرار نیستند و بنابراین مدل‌های رگرسیونی معمولی قادر به پاسخگویی در این شرایط نیستند. در این مواقع استفاده از مدل‌های رگرسیونی در محیط فازی می‌تواند راهگشای برخی از این مشکلات باشد.

رویکرد رگرسیون استوار یک مدل رگرسیون ایجاد می‌کند که تاثیرپذیری قوی نسبت به داده‌های پرت ندارد. نویسندگان زیادی به مطالعه و تحقیق بر روی مدل‌های رگرسیونی استوار در محیط فازی پرداخته‌اند. در زیر به برخی از آنها اشاره می‌گردد. با استفاده از رویکرد برنامه‌ریزی خطی فازی، پیتز (۱۹۹۴) رویکرد تاناکا و همکاران (۱۹۸۲) را به منظور کاهش اثر داده‌های پرت بهبود داده است. اوساش و اسکوتر (۲۰۰۲) با استفاده از حداقل مربعات پیراسته (LTS) و حداقل مربعات میانه (LMS) مدل رگرسیون فازی را ارائه کرده‌اند و عملکرد مدل پیشنهادی را در مواجهه با داده‌های پرت مورد مطالعه قرار داده‌اند. شون (۲۰۰۵) به منظور

^۱ نام ارائه دهنده مقاله : amir.khammar@birjand.ac.ir

اصلاح رویکرد تاناکا و همکاران (۱۹۸۲)، با استفاده از M -برآوردگرها یک روش استوارسازی برای مدل رگرسیون فازی پیشنهاد داده‌اند. چویی و باکلی (۲۰۰۸) مدل رگرسیون حداقل قدرمطلق انحرافات (LDA) فازی را معرفی کرده‌اند و نشان داده‌اند که مدل نسبت به داده‌های پرت استوار است. کولا و آپادین (۲۰۰۸) یک تحلیل رگرسیون فازی استوار براساس رتبه‌بندی مجموعه‌های فازی پیشنهاد کرده‌اند. دورسو و همکاران (۲۰۱۱) یک رگرسیون خطی فازی استوار براساس روش حداقل مربعات میانه-حداقل مربعات وزنی ($LMS - WLS$) را معرفی کرده‌اند. چاچی و همکاران (۲۰۱۶) با استفاده از روش کمترین قدرمطلقات (LA) دو مدل رگرسیون فازی چندگانه استوار ارائه کرده‌اند. چاچی و روزبه (۲۰۱۷) با استفاده از روش برآوردیابی حداقل مربعات پیراسته (LTS) یک برآورد استوار برای ضرایب مدل رگرسیونی فازی با داده‌های ورودی غیرفازی و داده‌های خروجی فازی معرفی کرده‌اند.

از برآوردگرهای استوار متداول می‌توان به M -برآوردگرها اشاره کرد که به راحتی قابل محاسبه بوده و خواص آماری متعددی دارند. در این مقاله با استفاده از M -برآوردگرها و فاصله علامتدار بین دو عدد فازی که در بخش بعد به آن پرداخته می‌شود، یک مدل رگرسیون خطی فازی استوار معرفی می‌کنیم. در ابتدا لازم است برخی مفاهیم و نکات درباره مجموعه‌های فازی بیان شود.

فرض کنید X یک مجموعه مرجع باشد. مجموعه فازی $\tilde{A}$ از X توسط تابع عضویت $\tilde{A}(x) : X \rightarrow [0, 1]$ تعریف می‌شود. α -برش مجموعه فازی $\tilde{A}$ ، به صورت $\tilde{A}[\alpha] = \{x | \tilde{A}(x) \geq \alpha\}$ برای $0 < \alpha \leq 1$ تعریف می‌شود.

تعریف ۱.۱. مجموعه $\tilde{A}$ از مجموعه مرجع X را یک عدد فازی گوئیم، اگر اولاً به ازای هر $x \in R$ ، $\tilde{A}(x) = 1$ و ثانیاً α -برش‌های $\tilde{A}$ به ازای $\alpha \in (0, 1]$ بازه‌هایی بسته‌همبند و کراندار باشند.

تعریف ۲.۱. فرض کنید ساختار تابع عضویت عدد فازی $\tilde{A}$ به صورت زیر باشد:

$$\mu_{\tilde{A}}(x) = \begin{cases} L\left(\frac{a-x}{l_a}\right), & x \leq a \\ R\left(\frac{x-a}{r_a}\right), & x > a \end{cases}$$

که در آن توابع شکل (مرجع) L و R ، توابعی غیرصعودی از R^+ به $[0, 1]$ هستند و $L(0) = R(0) = 1$. اگر $L(x) = R(x) = \max\{0, 1 - x\}$ ، آنگاه $\tilde{A}$ را یک عدد فازی مثلثی نامیده و با نماد $\tilde{A} = (a; l_a, r_a)_T$ نشان می‌دهیم. عدد حقیقی a مقدار نما (یا مرکز یا میانه) و اعداد مثبت l_a و r_a به ترتیب پهنای چپ و پهنای راست $\tilde{A}$ نامیده می‌شوند.

در صورتی که $L = R$ و $l_a = r_a$ ، آنگاه عدد فازی را متقارن نامیده و با نماد $\tilde{A} = (a; l)_T$ نشان می‌دهیم. برخی اعمال جبری برای اعداد فازی مثلثی بسیار ساده و دارای الگوهای مشخصی است و برای برخی دیگر از اعمال جبری نیز الگوهای تقریبی وجود دارد.

قضیه ۳.۱. اگر $\tilde{A} = (a; l_a, r_a)_T$ و $\tilde{B} = (b; l_b, r_b)_T$ دو عدد فازی مثلثی و $\lambda \in R$ باشد، آنگاه

$$\tilde{A} \oplus \tilde{B} = (a + b; l_a + l_b, r_a + r_b)_T,$$

$$\tilde{A} \ominus \tilde{B} = (a - b; l_a + r_b, r_a + l_b)_T,$$

$$\lambda \otimes \tilde{A} = \begin{cases} (\lambda a; \lambda l_a, \lambda r_a)_T, & \lambda > 0, \\ (\lambda a; -\lambda r_a, -\lambda l_a)_T, & \lambda < 0. \end{cases}$$

۲ کلاسی از فاصله‌های علامتدار بین دو عدد فازی

در این بخش با استفاده از اصل رتبه‌بندی اعداد فازی، کلاسی از فاصله‌های علامتدار بین دو عدد فازی را معرفی می‌کنیم. می‌گوییم رتبه عدد فازی $\tilde{A}$ کمتر از رتبه عدد فازی $\tilde{B}$ است ($\tilde{A} \prec \tilde{B}$) اگر و تنها اگر، $\tilde{A}$ قبل از $\tilde{B}$ قرار گرفته باشد. نویسندگان زیادی درباره رتبه‌بندی اعداد فازی تحقیق کرده‌اند که می‌توان به یاو و ویو (۲۰۰۰)، چن و هسی (۲۰۰۰)، کولا و آپادین (۲۰۰۸) و عارفی (۲۰۱۸) اشاره کرد.

تعریف ۱.۲. فرض کنید $\tilde{A}$ و $\tilde{B}$ دو عدد فازی و $d(\tilde{A}, \tilde{B})$ یک فاصله تعریف شده بین آنها باشد. تابع علامت زیر را در نظر بگیرید:

$$\gamma(\tilde{A}, \tilde{B}) = \begin{cases} +1, & \tilde{A} \succ \tilde{B}, \\ 0, & \tilde{A} \sim \tilde{B}, \\ -1, & \tilde{A} \prec \tilde{B}. \end{cases}$$

در این صورت کلاس فاصله‌های علامتدار بین $\tilde{A}$ و $\tilde{B}$ را به صورت زیر تعریف می‌کنیم:

$$D(\tilde{A}, \tilde{B}) = \gamma(\tilde{A}, \tilde{B}).d(\tilde{A}, \tilde{B}), \quad (۱)$$

در این مقاله، برای رتبه‌بندی اعداد فازی مثلثی از میانگین ارائه شده توسط یائو و ویو (۲۰۰۰) و فاصله پیشنهاد شده توسط عارفی (۲۰۱۸) استفاده می‌شود.

تعریف ۲.۲. (یائو و ویو (۲۰۰۰)) فرض کنید $\tilde{A}$ یک عدد فازی مثلثی با α -برش $A_\alpha = [A_\alpha^L, A_\alpha^R]$ باشد. میانگین $\tilde{A}$ به صورت زیر تعریف می‌شود:

$$I(\tilde{A}) = \frac{1}{\pi} \int_0^1 (L_\alpha^A + R_\alpha^A) d\alpha.$$

تعریف ۳.۲. (عارفی (۲۰۱۷)) برای دو عدد فازی مثلثی متقارن $\tilde{A} = (a; l_a)_L$ و $\tilde{B} = (b; l_b)_L$ داریم:

$$d(\tilde{A}, \tilde{B}) = |a - b| + \frac{1}{3} |l_a - l_b|.$$

با توجه به تعریف ۲.۲، برای دو عدد فازی $\tilde{A}$ و $\tilde{B}$ رتبه‌بندی به صورت زیر انجام می‌پذیرد:

$$\tilde{A} \succ \tilde{B} \Leftrightarrow I(\tilde{A}) > I(\tilde{B}), \quad \tilde{A} \sim \tilde{B} \Leftrightarrow I(\tilde{A}) = I(\tilde{B}), \quad \tilde{A} \prec \tilde{B} \Leftrightarrow I(\tilde{A}) < I(\tilde{B}).$$

با استفاده از فاصله تعریف شده در تعریف ۳.۲ و رتبه‌بندی فوق، (۱) فاصله علامتدار زیر را برای دو عدد فازی مثلثی متقارن $\tilde{A} = (a; l_a)_T$ و $\tilde{B} = (b; l_b)_T$ نتیجه می‌دهد:

$$D(\tilde{A}, \tilde{B}) = \begin{cases} |a - b| + \frac{1}{3} |l_a - l_b|, & I(\tilde{A}) > I(\tilde{B}), \\ 0, & I(\tilde{A}) = I(\tilde{B}), \\ -|a - b| - \frac{1}{3} |l_a - l_b|, & I(\tilde{A}) < I(\tilde{B}). \end{cases} \quad (۲)$$

با استفاده از ویژگی‌های فواصل فازی، قضیه زیر برای کلاس فاصله‌های علامتدار تعریف شده برقرار می‌باشد که از اثبات آن صرف نظر می‌کنیم.

لم ۴.۲. برای دو عدد فازی $\tilde{A}$ و $\tilde{B}$ ، ویژگی‌های زیر در مورد کلاس فاصله‌های علامتدار معرفی شده در تعریف ۱.۲ صدق می‌کند:

$$(۱) \text{ اگر } \tilde{A} \neq \tilde{B}, \text{ آنگاه } D(\tilde{A}, \tilde{B}) = -D(\tilde{B}, \tilde{A}).$$

$$(۲) \text{ اگر } \tilde{A} = \tilde{B}, \text{ آنگاه } D(\tilde{A}, \tilde{B}) = D(\tilde{B}, \tilde{A}).$$

$$(۳) \text{ اگر } \tilde{A} \succ \tilde{B} \succ \tilde{C}, \text{ آنگاه } D(\tilde{A}, \tilde{B}) + D(\tilde{B}, \tilde{C}) \geq D(\tilde{A}, \tilde{C}).$$

۳ رگرسیون خطی فازی استوار

فرض کنید می‌خواهیم پارامترهای مدل رگرسیون خطی فازی زیر را برآورد کنیم:

$$\tilde{Y}_i = a_0 \oplus (a_1 \otimes \tilde{x}_1) \oplus \dots \oplus (a_k \otimes \tilde{x}_n), \quad i = 1, 2, \dots, n, \quad (۳)$$

که در آن ضرایب رگرسیونی $a_0, \dots, a_k$ اعداد غیر فازی و داده‌های ورودی $\tilde{x}_i = (x_i, s_i)_T$ و داده‌های خروجی $\tilde{Y}_i = (y_i, z_i)_T$ ، $i = 1, 2, \dots, n$ ، اعداد فازی مثلثی متقارن می‌باشند. M -برآوردگر پارامتر مجهول a_j ، $j = 0, 1, \dots, k$ ، مقداری از a_j است که تابع $\sum \rho(\tilde{Y}_j - a_j \tilde{x}_j)$ را مینیمم کند. از آنجاییکه تفاوت دو عدد فازی عددی فازی است، ما قادر به مینیمم کردن مجموع مذکور در مدل رگرسیون فازی نخواهیم بود، اما فاصله بین دو عدد فازی می‌تواند عددی معمولی باشد. بنابراین در مدل رگرسیون فازی می‌توان از یک فاصله علامتدار بین مقادیر مشاهده شده و اعداد فازی برآورد شده استفاده کرد. با استفاده از حساب بین اعداد فازی مثلثی، معادله برآورد زیر را برای متغیر خروجی فازی برآورد شده $\hat{Y}$ خواهیم داشت:

$$\hat{Y}_i = a_0 \oplus (a_1 \otimes (x_1, s_1)_T) \oplus \dots \oplus (a_k \otimes (x_n, s_n)_T) = (a_0 + \sum_{i=1}^n a_i x_i, |a_0| + \sum_{i=1}^n |a_i| s_i)_T, \quad i = 1, 2, \dots, n, \quad (4)$$

برای بدست آوردن یک مدل بهینه و استوار از (4)، با استفاده از روش M -برآوردگرها، مجموع $\sum \rho(D(\tilde{Y}_i, \hat{Y}_i))$ را مینیمم کرده و محدودیت مربوط به تفاوت پهناها یعنی $|a_0| + \sum |a_i| s_i - z_i \leq c$ ، $i = 1, \dots, n$ ، را نیز لحاظ می‌کنیم. هدف از این محدودیت در نظر گرفته شده، تشخیص تفاوت بین پهنای خروجی z_i و پهنای برآورد شده $|a_0| + \sum |a_i| s_i$ و نگهداشتن این تفاوت کمتر از مقدار معین c است. برای بدست آوردن c مناسب از فرمول $c = 1/n \sum z_i$ ، $i = 1, 2, \dots, n$ ، استفاده می‌کنیم. لذا با حل مسئله مینیمم‌سازی زیر مواجه هستیم:

$$\begin{aligned} \min \quad & \sum_{i=1}^n \rho(D(\tilde{Y}_i, \hat{Y}_i)) \\ \text{s.t.} \quad & |a_0| + \sum_{i=1}^n |a_i| s_i - z_i \leq \frac{1}{n} \sum_{i=1}^n z_i, \quad i = 1, 2, \dots, n. \end{aligned} \quad (5)$$

در حالت خاص تابع $\rho(\cdot)$ را تابع چندکی در نظر می‌گیریم، بطوریکه به ازای $\tau \in (0, 1)$ ، $\rho(u) = u \cdot (\tau - I_{(u \leq 0)})$ ، که در آن $I(\cdot)$ تابع نشانگر می‌باشد. با قرار دادن تابع چندکی در (5)، مسئله مینیمم‌سازی زیر را خواهیم داشت

$$\begin{aligned} \min \quad & \sum_{i=1}^n D(\tilde{Y}_i, \hat{Y}_i) \cdot (\tau - I_{(D(\tilde{Y}_i, \hat{Y}_i) \leq 0)}) \\ \text{s.t.} \quad & |a_0| + \sum_{i=1}^n |a_i| s_i - z_i \leq \frac{1}{n} \sum_{i=1}^n z_i, \quad i = 1, 2, \dots, n. \end{aligned} \quad (6)$$

برای ساده‌سازی مسئله فوق، $D(\tilde{Y}_i, \hat{Y}_i)$ را با استفاده از متغیرهای نامنفی D_i^+ و D_i^- می‌توان به صورت زیر بازنویسی کرد:

$$D(\tilde{Y}_i, \hat{Y}_i) = \max\{d(\tilde{Y}_i, \hat{Y}_i), 0\} I_{(y_i > a_0 - \sum_{i=1}^n a_i x_i)} - \max\{-d(\tilde{Y}_i, \hat{Y}_i), 0\} I_{(y_i < a_0 - \sum_{i=1}^n a_i x_i)} = D_i^+ - D_i^-.$$

از طرفی، با توجه به تعریف ۳.۲ و با استفاده از متغیرهای نامنفی d_i^+ و d_i^- داریم:

$$d(\tilde{Y}_i, \hat{Y}_i) = (y_i - a_0 - \sum_{i=1}^n a_i x_i) + \frac{1}{\varphi} (z_i - |a_0| - \sum_{i=1}^n |a_i| s_i) = d_i^+ - d_i^-.$$

در رابطه فوق، اگر $(y_i - a_0 - \sum_{i=1}^n a_i x_i) + \frac{1}{\varphi} (z_i - |a_0| - \sum_{i=1}^n |a_i| s_i) \geq 0$ آنگاه $d_i^+ = (y_i - a_0 - \sum_{i=1}^n a_i x_i) + \frac{1}{\varphi} (z_i - |a_0| - \sum_{i=1}^n |a_i| s_i)$ و $d_i^- = 0$ و اگر $(y_i - a_0 - \sum_{i=1}^n a_i x_i) + \frac{1}{\varphi} (z_i - |a_0| - \sum_{i=1}^n |a_i| s_i) < 0$ آنگاه $d_i^+ = 0$ و $d_i^- = -(y_i - a_0 - \sum_{i=1}^n a_i x_i) - \frac{1}{\varphi} (z_i - |a_0| - \sum_{i=1}^n |a_i| s_i)$.

$|a_{\circ}| - \sum_{i=1}^n |a_i|s_i$ در نهایت مسئله مینیم سازی به مسئله برنامه ریزی زیر تبدیل می شود:

$$\begin{aligned}
 & \min \sum_{i=1}^n (D_i^+ \tau + D_i^- (1 - \tau)) \\
 & s.t. D_i^+ - D_i^- = \max\{d_i^+ - d_i^-, \circ\} I_{(y_i \geq a_{\circ} - \sum_{i=1}^n a_i x_i)} - \max\{-(d_i^+ - d_i^-), \circ\} I_{(y_i < a_{\circ} - \sum_{i=1}^n a_i x_i)} \\
 & d_i^+ - d_i^- = (y_i - a_{\circ} - \sum_{i=1}^n a_i x_i) + \frac{1}{3} (z_i - |a_{\circ}| - \sum_{i=1}^n |a_i|s_i) \\
 & |a_{\circ}| + \sum_{i=1}^n |a_i|s_i - z_i \leq \frac{1}{n} \sum_{i=1}^n z_i, \\
 & D_i^+, D_i^-, d_i^+, d_i^- \geq \circ; \quad i = 1, 2, \dots, n. \\
 & a_j \in R, \quad j = \circ, 1, \dots, n.
 \end{aligned} \tag{7}$$

با حل مسئله برنامه ریزی فوق، مدل رگرسیون چندکی فازی استوار نتیجه می شود که به اختصار آن را مدل RFQR می نامیم. لازم به ذکر است که در این مقاله برای حل مسائل برنامه ریزی از نرم افزار لینگو استفاده می شود.

۱.۳ مثال عددی

در این بخش با استفاده از یک مثال عددی، مقایسه ای بین مدل جدید RFQR و مدل شناخته شده حداقل مربعات (LS) در محیط فازی ارائه می شود. در روش LS در مدل رگرسیونی فازی، به دنبال مینیم سازی مجموع مربعات فواصل بین مقادیر مشاهده شده و اعداد فازی برآورد شده، یعنی $\min \sum d^2(\tilde{Y}_i, \hat{Y}_i)$ هستیم. در اینجا از نماد توان دوم به جای نماد قدرمطلق در فاصله تعریف شده در تعریف ۳.۲ استفاده می کنیم، بطوریکه خواهیم داشت:

$$d^2(\tilde{Y}_i, \hat{Y}_i) = (y_i - a_{\circ} - \sum_{i=1}^n a_i x_i)^2 + \frac{1}{3} (z_i - |a_{\circ}| - \sum_{i=1}^n |a_i|s_i)^2. \tag{8}$$

داده های موجود در جدول ۱ را در نظر بگیرید. این مجموعه داده که اولین بار توسط ساکاو و یانو (۱۹۹۲) ارائه شده است توسط محققان مختلفی مورد بررسی قرار گرفته است. مدل رگرسیونی در اینجا همانند مدل (۳) با اعداد فازی مثلثی مقارن می باشد. در ادامه مقایسه ای بین مدل پیشنهاد شده (مدل RFQR) و مدل شناخته شده حداقل مربعات خطا (LS) ارائه می شود. با استفاده از مدل LS بر مبنای فاصله (۸)، مدل رگرسیونی زیر حاصل می شود (شکل ۱):

جدول ۱: داده های مثال ۱.۳.

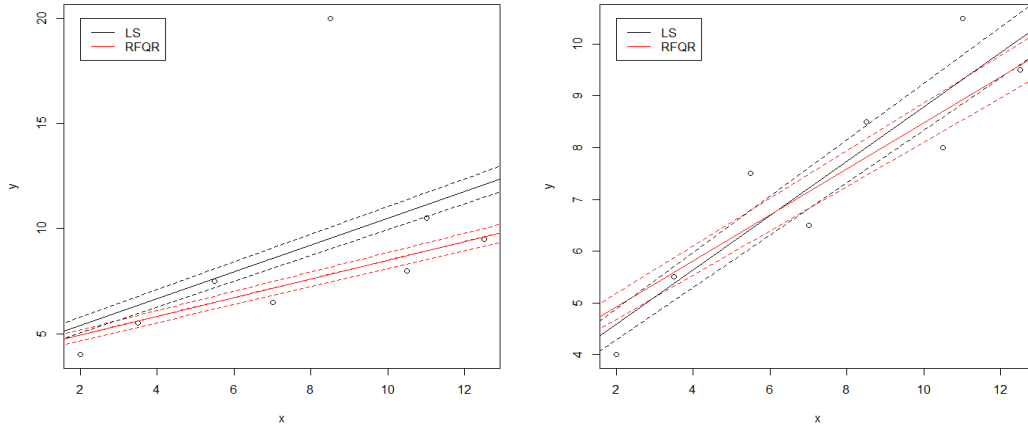
شماره	$\tilde{Y}$	$\tilde{X}$
۱	(۲, ۵)	(۴, ۵)
۲	(۵, ۵, ۵)	(۳, ۵)
۳	(۷, ۵, ۱)	(۵, ۵, ۵)
۴	(۶, ۵, ۵)	(۷, ۵)
۵	(۸, ۵, ۵)	(۸, ۵, ۵)
۶	(۸, ۱)	(۱۰, ۵, ۱)
۷	(۱۰, ۵, ۵)	(۱۱, ۵)
۸	(۹, ۵, ۵)	(۱۲, ۵, ۵)

$$\hat{Y}_{LS} = ۳/۵۲۹ \oplus \circ/۵۲۵ \otimes \tilde{X}.$$

مدل RFQR ارائه شده در بخش ۵، برای $\tau = ۰/۵$ مدل رگرسیونی زیر را بدست می‌دهد (شکل ۱):

$$\hat{Y}_{RFQR} = ۴/۰۳۷ \oplus ۰/۴۴۴ \otimes \tilde{X}. \quad (۹)$$

نمودار مرتبط با مدل‌های فوق در شکل ۱ نشان داده شده است. حال بجای مقدار پاسخ مشاهده پنجم مقدار $\tilde{y}_5 = (۲۰, ۰/۵)_T$ را که داده‌ای پرت است جایگزین می‌کنیم. همانگونه که در شکل ۱ مشخص است برخلاف مدل LS، مدل RFQR تحت تاثیر داده پرت قرار نگرفته است.



شکل ۱: مدل‌های رگرسیونی مثال ۵ با کران‌های بالا و پایین.

۴ نیکویی برازش

در این بخش، با استفاده از یک شاخص نیکویی برازش تعریف شده به وسیله چن و هسوئه (۲۰۰۷)، یک مطالعه مقایسه‌ای بین مدل جدید RFQR و برخی مدل‌های شناخته شده دیگر ارائه می‌شود. چن و هسوئه (۲۰۰۷)، معیار فاصله براساس α -برش‌ها را به عنوان یک شاخص نیکویی برازش در مدل‌های رگرسیون فازی به صورت زیر تعریف کرده‌اند.

تعریف ۱۰۴. برای مدل رگرسیون خطی فازی (۳)، فرض کنید D_i میانگین خطا برای i امین مورد توسط m α -برش باشد، بطوریکه

$$D_i = \frac{1}{m} \sum_{k=1}^m (|(\hat{Y}_i)_{\alpha_k}^L - (\tilde{Y}_i)_{\alpha_k}^L| + |(\hat{Y}_i)_{\alpha_k}^U - (\tilde{Y}_i)_{\alpha_k}^U|),$$

که برای یک α -برش مشخص، $(\hat{Y}_i)_{\alpha_k}^L$ و $(\hat{Y}_i)_{\alpha_k}^U$ به ترتیب کران‌های پایین و بالای i امین متغیر خروجی فازی برآورد شده و $(\tilde{Y}_i)_{\alpha_k}^L$ و $(\tilde{Y}_i)_{\alpha_k}^U$ کران‌های پایین و بالای i امین متغیر خروجی فازی مشاهده شده می‌باشند.

متوسط فواصل مذکور (شاخص MD) را می‌توان به عنوان شاخص نیکویی برازش برای مدل (۳) در نظر گرفت:

$$MD = \frac{1}{n} \sum_{i=1}^n D_i.$$

چن و هسوئه (۲۰۰۷) از دو α -برش ۰ و ۱ برای مقایسه برخی مدل‌های فازی استفاده کرده‌اند. ما نیز در مثال بعد از این دو α -برش برای مقایسه مدل‌های مورد نظرم‌ان استفاده می‌کنیم. مجموعه داده‌های جدول ۱ را در نظر بگیرید. در این مثال با حضور و عدم حضور داده پرت رویکرد پیشنهادی را با رویکردهای چن و هسوئه

(۲۰۰۹)، هاشمپور و همکاران (۲۰۰۹) و طاهری و کلکینما (۲۰۱۲) مقایسه کرده‌ایم. مدل‌های رگرسیونی به صورت زیر بدست می‌آید. مقادیر MD آنها در جدول ۲ قابل مشاهده می‌باشد.

$$\hat{Y}_{CH} = ۳/۵۷۵۰ \oplus ۰/۵۱۹۶\tilde{X} \oplus (۰, ۰/۳۰۰۹)_T, \quad \hat{Y}_{HMY} = ۳/۹۴۴۴ \oplus ۰/۴۴۴۴\tilde{X},$$

$$\hat{Y}_{TK} = ۳/۹۴۴۴ \oplus ۰/۴۴۴۴\tilde{X} \oplus (۰, ۰/۲۷۷۸)_T, \quad \hat{Y}_{RFGR} = ۴/۰۳۷۰ \oplus ۰/۴۴۴۴\tilde{X}.$$

با توجه به اطلاعات جدول ۲، براساس معیار MD مدل پیشنهادی RFQR عملکرد را نشان می‌دهد. علاوه بر این، در حضور داده پرت $\tilde{y}_5 = (۱۸/۵, ۰/۵)_T$

جدول ۲: مقادیر MD برای داده‌های جدول ۱.

شماره	CH	HMY	TK	RFQR _{T=۰.۵}
۱	۰,۶۱۳۰	۰,۸۳۳۲	۰,۸۳۳۲	۰,۹۲۵۰
۲	۰,۱۰۸۵	۰,۱۳۹۰	۰,۰۰۰۲	۰,۰۹۳۰
۳	۱,۰۷۰۵	۱,۱۱۱۴	۱,۱۱۱۴	۱,۰۲۱۰
۴	۰,۷۰۸۰	۰,۵۵۵۲	۰,۵۵۵۲	۰,۶۴۵۰
۵	۰,۵۱۳۵	۰,۷۷۸۲	۰,۷۷۸۲	۰,۶۸۹۰
۶	۱,۰۲۴۵	۰,۶۱۰۶	۰,۶۱۰۶	۰,۶۹۹۰
۷	۱,۲۱۶۰	۱,۶۶۷۲	۱,۶۶۷۲	۱,۳۲۹۰
۸	۰,۵۶۲۵	۰,۱۳۹۲	۰,۰۰۰۶	۰,۰۹۵۵
(MD) بدون داده پرت	۰,۷۲۷۱	۰,۷۲۹۲	۰,۶۹۴۶	۰,۶۸۷۰
(MD) با داده پرت	۲,۳۰۶۰	۱,۹۷۹۲	۱,۹۴۴۶	۱,۹۳۷۰

نیز مدل RFQR عملکرد بهتری را نشان می‌دهد.

نتیجه‌گیری

در این مقاله یک روش جدید استوارسازی برای برآورد ضرایب در مدل رگرسیونی فازی براساس M -برآوردگرها و یک فاصله علامتدار ارائه شده است. مدل رگرسیونی خطی فازی در نظر گرفته شده در این مقاله، شامل داده‌های ورودی و خروجی فازی مثلثی متقارن و ضرایب رگرسیونی غیرفازی می‌باشد. روش پیشنهاد شده در حالت خاص مدل رگرسیون چندکی فازی استوار را نتیجه می‌دهد. مثال عددی و شاخص نیکویی برازش نشان می‌دهند که مدل رگرسیون چندکی فازی استوار ارائه شده، نسبت به برخی از مدل‌های رگرسیونی فازی شناخته شده، بخصوص در حالت وجود داده پرت در مجموعه داده‌ها عملکرد بهتری دارد.

مراجع

- [16] Aefi M. (2018), Testing statistical hypotheses under fuzzy data and based on a new signed distance, *Iranian J. Fuzzy Syst.*, in Press, DOI: 10.22111/IJFS.2017.3361.
- [16] Chachi J., Taheri S.M., Fattahi S., and Hosseini Ravandi S.A. (2016), Two Robust Fuzzy Regression Models and Their Applications in Predicting Imperfections of Cotton Yarn, *Journal of Textiles and Polymers*, 4, 60-68.
- [16] Chachi J. and Roozbeh M. (2017), A Fuzzy Robust Regression Approach Applied to Bedload Transport Data, *Communications in Statistics-Simulation and Computation*, 46, 1703-1714.

- [16] Chang P.T. and Lee C.H. (1994), Fuzzy least absolute deviations regression based on the ranking of fuzzy numbers, *In: Proc. of the Third IEEE World Congress on Computational Intelligence, Orlando, FL*, 2, 1365-1369.
- [5] Chen L.H. and Hsueh C.C. (2007), A mathematical programming method for formulating a fuzzy regression model based on distance criterion, *IEEE Transactions on Systems, Man, Cybernetics B*, 37, 705-712.
- [6] Chen L.H. and Hsueh C.C. (2009), Fuzzy regression models using the least-squares method based on the concept of distance, *IEEE Transactions on Fuzzy Systems*, 17, 1259-1272.
- [7] Choi S.H. and Buckley J.J. (2008), Fuzzy regression using least absolute deviation estimators, *Soft Computing*, 12, 257-263.
- [16] D'Urso P., Massari R. and Santoro A. (2011), Robust Fuzzy Regression Analysis. *Information Sciences*, 181, 4154-4174.
- [9] Hassanpour H., Maleki H.R. and Yaghoobi M.A. (2009), A goal programming approach to fuzzy linear regression with non-fuzzy input and fuzzy output data, *Asia-Pacific Journal of Operational Research*, 26, 587-604.
- [16] Kula K.S. and Apaydin A. (2008), Fuzzy Robust Regression Analysis Based on the Ranking of Fuzzy Sets, *International Journal Uncertainty, Fuzziness and Knowledge-Based Systems*, 16, 663-681.
- [16] Oussalah, M. and De Schutter J. (2002), Robust Fuzzy Linear Regression and Application for Contact Identification, *Intelligent Automation and Soft Computing*, 8, 31-39.
- [16] Peters G. (1994), Fuzzy linear regression with Fuzzy intervals, *Fuzzy Sets and Systems*, 63, 45-55.
- [16] Sakawa M. and Yano H. (1992), Multiobjective fuzzy linear regression analysis for fuzzy input-output data, *Fuzzy Sets and Systems*, 47, 173-181.
- [16] Shon B.Y. (2005), Robust Fuzzy Linear Regression Based on M-estimators, *J. Appl. Math. & Computing*, 18, 591-601.
- [16] Taheri S.M. and Kelkinnama M. (2012), Fuzzy linear regression based on least absolute deviations, *Iranian J. Fuzzy Syst.*, 9, 121-140.
- [16] Tanaka H., Uejima S., Asai K. and Linea (1982), Regression Analysis with Fuzzy Model, *IEEE Trans. Systems, Man, and Cybernetics*, 12, 903-907.

نمودارهای کنترل فازی براساس پروفایل‌ها

مریم جلیلیان عاملی^۱، بهرام صادق‌پور گیلده^۱، زینب عباسی گنجی^۱

^۱ دانشگاه فردوسی مشهد، دانشکده علوم ریاضی، گروه آمار

^۲ مرکز تحقیقات جهاد کشاورزی مشهد

چکیده

چکیده: به مجموعه عملیاتی مانند اندازه‌گیری یا آزمون که روی یک محصول یا کالا انجام می‌شود تا مشخص شود آیا آن محصول با مشخصات فنی مورد نظر مطابقت دارد یا خیر کنترل کیفیت گویند که نقش مهمی در افزایش کیفیت محصول ایفا می‌کند. یکی از ساده‌ترین روش‌های کنترل آماری فرآیند حین تولید نمودار کنترل است. در صورتی که مشاهدات به صورت زبانی یا نادقیق باشند استفاده از نظریه مجموعه‌های فازی و به تبع آن نمودارهای کنترل فازی نسبت به نمودار کنترل کلاسیک (شوهارت) حساس‌تر هستند. لذا استفاده صحیح از نمودار کنترل فازی منجر به تولید محصولات با کیفیت می‌شود. این حوزه بسیار پیچیده است زیرا شامل حوزه وسیعی از صنایع است. در این تحقیق، مروری کلی از نمودار کنترل، تأثیر پروفایل‌ها بر نمودار کنترل، نمودار کنترل فازی و نمودار کنترل فازی بر اساس پروفایل‌ها ارائه می‌دهیم.

کلمات کلیدی: نمودار کنترل، پروفایل، پایش پروفایل، نمودار کنترل فازی

۱ مقدمه

کنترل کیفیت نقش مهمی در افزایش کیفیت محصول ایفا می‌کند و ^۱کنترل کیفیت آماری یکی از شاخه‌های کنترل کیفیت می‌باشد. کنترل کیفیت آماری شامل مباحثی نظیر نمونه‌گیری به منظور پذیرش، کنترل فرآیند آماری ^۲SPC و ^۳طراحی آزمایش‌ها می‌باشد. بهبود کیفیت را می‌توان کاهش در تغییرپذیری فرآیند و محصول تعریف نمود. کنترل فرآیند آماری یک روش آماری است که برای کاهش پراکندگی و در نتیجه بهبود کیفیت استفاده می‌شود. هدف اصلی کنترل فرآیند آماری حذف تغییرپذیری فرآیند است. در کاربردهای کنترل فرآیند آماری فرض می‌شود که کیفیت یک فرآیند یا محصول را می‌توان به طور مناسب با توزیع یک ویژگی کیفی توصیف کرد. با این حال، در موقعیت‌های عملی متعدد، کیفیت یک فرآیند یا محصول به طور خلاصه با رابطه بین متغیر پاسخ کیفیت و یک یا چند متغیر توضیحی ارائه می‌شود. این رابطه

^۱مریم جلیلیان عاملی : ameli11365@yahoo.com

^۱Statistical Quality Control

^۲Statistical Process Control

^۳Design Of Experiment

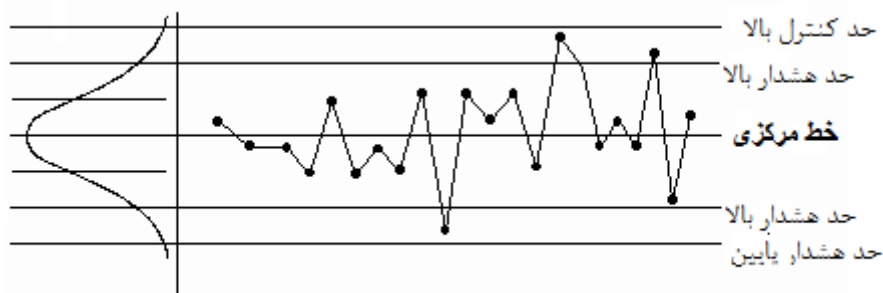
به عنوان^۴ پروفایل شناخته می‌شود. بنابراین پایش و کنترل این نوع فرآیندها و محصولات با پایش بر پروفایل‌ها انجام می‌شود. به عبارت دیگر، پایش پروفایل ایجاد^۵ نمودارهای کنترل برای مواردی است که در آن کیفیت یک فرآیند یا محصول می‌تواند با یک رابطه کارکردی بین متغیر پاسخ و یک یا چند متغیر توضیحی شناخته شود. (۱)

۲ نمودارهای کنترل

یکی از ابزارهای قدرتمند کنترل فرآیند آماری، نمودارهای کنترل می‌باشند. نمودار کنترل جهت کنترل میزان تغییرات در یک مشخصه کیفی یا چندین مشخصه کیفی مورد استفاده قرار می‌گیرد. در دهه (۱۹۲۰)، استوارت اولین بار استفاده از نمودارهای کنترل را مطرح کرد. نمودار کنترل زمانی استفاده می‌شود که بخواهیم بر تغییرات یک فرآیند تولید نظارت کنیم. این نمودارها مشخص می‌کنند که آیا یک فرآیند تحت کنترل است یا خیر؟ سه کاربرد عمده این نمودارها عبارتند از:

- کاهش تغییر پذیری فرآیند
- نظارت و کنترل یک فرآیند
- تخمین پارامترهای محصول یا فرآیند.

نمودار کنترل مشخصه کیفی را که بر اساس اطلاعات نمونه اندازه‌گیری یا محاسبه شده است بر حسب شماره نمونه یا زمان نشان می‌دهد. نمودار کنترل شامل یک خط مرکزی^۶ CL است که مقدار متوسط مشخصه کیفی را در حالت تحت کنترل نشان می‌دهد و یا به عبارت دیگر حالتی از فرآیند را نشان می‌دهد که فقط خطاهای تصادفی حضور دارند. چهار خط افقی دیگر که در نمودار کنترل قرار دارند. حد کنترل بالا^۷ UCL حد هشدار بالا^۸ UWL حد هشدار پایین^۹ LWL و حد کنترل پایین^{۱۰} LCL نامیده می‌شود. اگر فرآیند تحت کنترل باشد کلیه نقاطی که به منظور ارزیابی و بر اساس اطلاعات نمونه محاسبه شده اند بین این حدود واقع می‌شوند. (۲)



شکل ۱: یک نمودار کنترل

۳ نقش پروفایلها در نمودارهای کنترل

همان‌طور که ذکر شد رابطه بین یک متغیر پاسخ و یک یا چند متغیر توضیحی پروفایل نامیده می‌شود. این رابطه باید در طول زمان کنترل شود. با توجه به رابطه رگرسیونی بین متغیر پاسخ و متغیر توضیحی، انواع پروفایل‌های خطی ساده، چند متغیره، غیرخطی و ... وجود دارد. پایش پروفایل اساساً بررسی ثبات یک رابطه رگرسیونی در

^۴Profile

^۵Control Chart

^۶Center Line

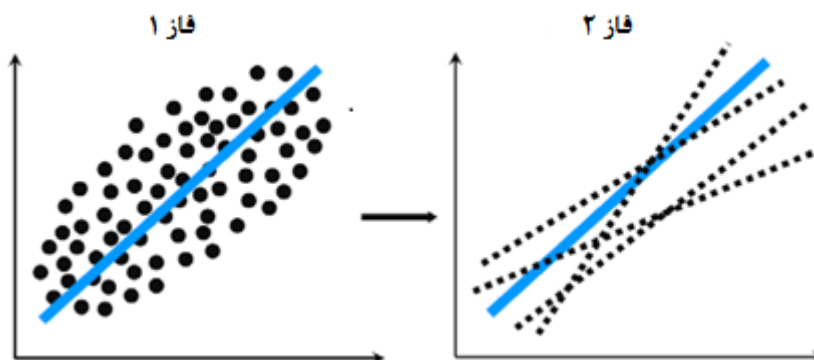
^۷Upper Control Limit

^۸ Upper Warning Limit

^۹ Lower Warning Limit

^{۱۰}Lower Control Limit

طی زمان بر اساس داده‌های مشاهده شده است. ساخت نمودارهای کنترل برای تشخیص اینکه آیا این رابطه در طول زمان تغییر کرده‌است یا نه، به عنوان پایش پروفایل شناخته شده‌است، که اخیراً در کنترل کیفیت توسعه‌یافته و رایج شده‌است. پایش پروفایل می‌تواند در دو فاز مورد مطالعه قرار گیرد. فاز ۱ و فاز ۲. بین این دو فاز از لحاظ اهداف و نحوه عملکرد تفاوت‌هایی وجود دارد و باید هنگام پایش فرایندها در فاز ۱ و ۲ دقت کافی به عمل آورد. در فاز ۱ با مجموعه‌ای از داده‌ها که از گذشته فرآیند در دسترس است سروکار داریم. اهداف فاز ۱ شامل به دست آوردن اطلاعات در مورد پراکندگی فرآیند در طول زمان، ارزیابی پایداری فرآیند، و مدل کردن عملکرد فرآیند تحت کنترل است. در فاز ۱، مجموعه‌ای از داده‌های فرآیند جمع‌آوری می‌شود و مورد تجزیه و تحلیل قرار می‌گیرند و بر اساس آنها حدود کنترلی آزمایشی طراحی می‌شود. سپس با استفاده از این حدود آزمایشی به دست آمده تعیین می‌شود که آیا فرآیند در طی این دوره زمانی در حالت تحت کنترل قرار داشته است یا خیر؟ دیگر این که آیا این حدود کنترلی حدود قابل اطمینانی برای پایش داده‌های آینده‌ی این فرآیند هستند؟ در حقیقت هدف فاز ۱، اطمینان از ثبات و پایداری فرآیند و برآورد پارامترها و حدود کنترلی برای پایش داده‌های آینده در فاز ۲ است. معیار ارزیابی نمودارهای کنترل در فاز ۱، احتمال کشف انحرافات با دلیل است. در این فاز، ابتدا مقدار ثابتی را برای احتمال خطای نوع یک در نظر گرفته و براساس آن حدود کنترل محاسبه می‌شود. سپس برای مقایسه‌ی رویکردهای مختلف در این فاز، با اعمال انحرافات بادلیل مختلف در پارامترهای فرآیند، احتمال کشف این احتمالات با دلیل برای نمودارهای کنترل متفاوت، مقایسه می‌شود. در فاز ۲ معمولاً فرض می‌شود که فرآیند در ابتدا از ثبات لازم برخوردار است و پارامترهای آن مشخص هستند و یا در فاز ۱ برآورد شده‌اند. در فاز ۲ از خط مرجعی که در فاز ۱ ایجاد شده‌است استفاده می‌شود. نگرانی اصلی در تجزیه و تحلیل فاز ۲، تشخیص سریع تغییرات در پارامترهای فرآیند از طریق پایش آنلاین بر پروفایل فرآیند است.



شکل ۲: فاز ۱ و فاز ۲ در پایش پروفایل

۴ نمودار کنترل فازی

در پایش پروفایل‌ها با ویژگی‌های کیفی دقیق که در محیط غیرفازی انجام می‌شود، فرض بر این است که داده‌های ورودی و خروجی فازی نیستند. اما در بسیاری از کاربردهای واقعی، داده‌ها دارای نوعی عدم قطعیت، نوسانات، عدم دقت و ابهام هستند و یا ممکن است داده‌های زبانی باشند. اطلاعات ممکن است ناقص، (غیردقیق، مبهم) باشند، و هر یک از این نقص اطلاعات منجر به انواع مختلف عدم قطعیت می‌شود. نظریه مجموعه‌های فازی برای اولین بار توسط (۱۸) معرفی شد. این نظریه برای توضیح عدم قطعیت و عدم دقت استفاده می‌شود. در برخی از کاربردهای واقعی، یک مدل رگرسیون خطی فازی می‌تواند پروفایل مناسبی را نشان دهد که در آن ویژگی کیفیت متغیر پاسخ فازی است. برخلاف روش‌های سنتی در پایش پروفایل‌ها، رویکردهای فازی می‌تواند بدون از دست دادن عملکرد و اثربخشی، با عدم دقت و عدم اطمینان مواجه شوند. بنابراین پایش پروفایل‌های فازی می‌تواند برای طیف وسیعی از کاربردها به‌کار رود. با به‌کارگیری نظریه مجموعه‌های فازی در کنترل فرایندها می‌توانیم کیفیت محصولات را دقیق‌تر کنترل کنیم. به همین دلیل، نمودارهای کنترل فازی نسبت به نمودار کنترل شوهارت حساس‌تر هستند و استفاده از نمودارهای کنترل فازی منجر به تولید محصولات با کیفیت می‌شود. (۱۲) برای اولین بار از مجموعه‌های فازی به عنوان مبانی برای توضیح اندازه‌گیری انطباق هر واحد محصول با مشخصات فنی تعیین شده استفاده کرد. بسیاری از پژوهشگران روش‌های مختلفی را برای ساخت نمودارهای کنترل بر اساس داده‌های فازی معرفی کرده‌اند.

(برای اطلاع بیشتر به (۱۳) مراجعه شود). در این مقاله، برخی رویکردهای پایش پروفایل‌های خطی فازی را معرفی می‌کنیم که در آن رابطه بین متغیر پاسخ فازی و برخی از متغیرهای توضیحی با یک رابطه رگرسیون خطی تعریف می‌شود. نمودارهای کنترل فازی را بر اساس مشخصه وصفی و متغیر طبقه‌بندی می‌کنیم. نمودار کنترل فازی وصفی به نمودارهای U, NP, C, P معروف‌اند. نمودار کنترل فازی متغیرها را می‌توان به نمودارهای میانگین دامنه $(\bar{X} - R)$ ، میانگین انحراف معیار $(\bar{X} - S)$ ، واریانس (S^2) ، میانگین دامنه متحرک $(\bar{X} - MR)$ نمودار کنترل جمع تجمعی $CUSUM^{11}$ ، میانگین متحرک موزون نمایی $EWMA^{12}$ ، T^2 هتلینگ شناخته شده‌اند.

۵ پایش پروفایل‌های فازی

پایش پروفایل‌های فازی توسط (۶) مورد بررسی قرار گرفت. آن‌ها یک روش تک متغیری برای پایش پروفایل‌های کیفیت فازی را ایجاد کردند و یک برنامه‌ریزی خطی چند هدفه برای ساخت پروفایل‌های فازی را طراحی نمودند. آن‌ها از نمودارهای کنترل فازی برای پایش شیب و عرض از مبدا پروفایل‌ها و به طور جداگانه استفاده کردند. (۳)

توان روش پایش تک متغیره را برحسب احتمال هشدار محاسبه کرده و استفاده از رویکرد خود برای شناسایی تغییرات متوسط و بزرگ در پروفایل فرآیند را توصیه کردند. در تحقیقات قبلی در مورد پایش پروفایل، اکثر مطالعات انجام‌شده در پایش پروفایل‌ها با ویژگی‌های کیفی دقیق در محیط غیر فازی انجام شد. با این حال، در بسیاری از کاربردهای واقعی، این داده‌ها دارای نوعی عدم قطعیت، نوسانات، و ابهام هستند و یا ممکن است داده‌های زبانی باشند. تاکنون کارهای اندکی برای پایش پروفایل‌های فازی در هر دو فاز ۱ و ۲ انجام شده‌است. اگر کسی به تشخیص تغییرات پایدار کوچک و متوسط علاقمند باشد، انواع دیگر جدول‌های کنترل ممکن است ترجیح داده شوند؛ به خصوص، میانگین متحرک موزون نمایی چندمتغیره $MEWMA^{13}$ را می‌توان مورد استفاده قرار داد. نمودار کنترل $MEWMA$ مشابه نسخه تک متغیره آن، نسبت به تغییرات کوچک و متوسط حساس است بنابراین، می‌تواند یک ابزار موثر برای پایش برخی از فرآیندهای حساس یا پروفایل‌های محصول باشد. دو رویکرد اصلی در بررسی مقالات مربوط به پایش پروفایل به منظور دستیابی به یک فرآیند پایدار و پروفایل پایدار وجود دارد. ۱- رویکرد تک متغیری که از نمودارهای کنترل تک متغیری استفاده می‌کند و ۲- رویکرد چند متغیره که از نمودارهای کنترل چند متغیره استفاده می‌کند. (۴)

نشان دادند که نمودارهای کنترل چند متغیره، نسبت به نمودارهای کنترل تک متغیره، حساس‌تر هستند. انواع مختلفی از نمودارهای کنترل چند متغیره، از جمله T^2 - هتلینگ، نمودار کنترل جمع تجمعی چندمتغیره $MCUSUM^{14}$ و $MEWMA$ در تلاش برای بهبود پایش همزمان ویژگی‌های کیفیت توسعه یافته‌اند. نمودارهای کنترل جمع تجمعی چندمتغیره و میانگین متحرک موزون نمایی چندمتغیره از اطلاعات فرآیند هر دو نمونه قبلی و فعلی استفاده می‌کنند. بنابراین، آن‌ها در مقایسه با نمودار کنترل T^2 که تنها از اطلاعات فعلی فرآیند بهره می‌برند، نسبت به تغییرات کوچک فرآیند حساس‌تر هستند. بنابراین، یکی از این نمودارهای کنترل چند متغیره می‌تواند براساس علاقه محقق با توجه به سرعت در تشخیص سریع تغییرات بزرگ یا کوچک در فرآیند انتخاب شود. در پایش فرآیندها بر اساس پروفایل‌ها دو روش چند متغیره اصلی متداول است. روش اول که توسط (۹) بکار گرفته شده به شرح زیر است: ۱) ساخت پروفایل برای هر نمونه (۲) پایش ضرایب پروفایل با استفاده از نمودارهای کنترل تک متغیری یا چند متغیره، و ۳) محاسبه ضرایب پروفایل فرآیند پایدار با میانگین‌گیری پارامترهای پروفایل به طور متوسط. روش دوم روش توسط مستک و همکاران (۱۹۹۴) به شرح زیر است: ۱) نظارت بر ویژگی کیفیت به عنوان متغیر پاسخ با استفاده از نمودار کنترل چند متغیره، ۲) ایجاد مدل پروفایل برای هر نمونه پس از دستیابی به فرآیند تحت کنترل، و ۳) محاسبه ضرایب پروفایل فرآیند پایدار از طریق میانگین‌گیری مجموعه نمونه تحت کنترل.

¹¹Cumulative Sum(or CUSUM) Control Chart

¹²Exponentially Weighted Moving Average

¹³Multivariate Exponentially Weighted Moving Average

¹⁴Multivariate Cumulative Sum(or MCUSUM) Control Chart

۶ نمودارهای کنترل فازی بر اساس پروفایلها

اولین مقاله در این زمینه توسط (۶) نوشته شده در این کار نمودارهای کنترل انفرادی ترسیم شده است. در واقع توسعه نمودار دامنه متحرک $I-MR^{15}$ برای مشاهدات فازی است. یک مدل رگرسیون خطی فازی برای ساخت پروفایلها با استفاده از برنامه نویسی خطی ارائه شده است سپس نمودار $I-MR$ را بر اساس مشاهدات تکی فازی جهت پایش شیب پروفایل توسعه دادند مطالعه ای در مورد رضایت مشتری برای نشان دادن کاربرد روش و بیان تحلیل حساسیت پارامترها برای ایجاد یک پروفایل فازی ارائه شده است. تحلیل پروفایلها شامل دو مرحله است فاز ۱ و فاز ۲. قبل از کار نقدریان و قبادی (۲۰۱۲) مطالعه ای برای پایش پروفایلهای فازی در هر دو مرحله ای فاز ۱ و فاز ۲ انجام نشده بود. در این مقاله برخی روشهای تک متغیری برای پایش پروفایلهای خطی فازی را معرفی می کنیم. در تحلیل فازی سعی بر آن است که هر کدام از خطوط رگرسیون فازی خارج از کنترلی را بیاییم و آنها را از مجموعه داده ها حذف کنیم با فرض اینکه عوامل اسنادپذیر متناظر را می توان شناسایی و حذف کرد. این برای تخمین دقیق پارامترهای رگرسیون فازی تحت کنترل برای استفاده در فاز ۲ لازم است. در ابتدا مقادیر x را کد می کنیم بطوریکه میانگین مقدار کد شده صفر است. این عمل تحلیل را آسان تر می کند و موجب حذف روشهای پایش چند متغیره می شود. چهار روش برای غیرفازی سازی وجود دارد ۱- میان دامنه α برش ۲- میانه α برش ۳- میانگین α برش ۴- مد α برش در این مقاله میان دامنه α برش به عنوان روش غیر فازی ساز به کار گرفته شده است. مشاهدات انفرادی فازی و دامنه متحرک $I-MR$ برای بررسی پایداری رابطهی کارکردی فازی و تخمین پارامترهای فرآیند کنترل توسعه داده شده اند. حال به بررسی مقاله دوم در این زمینه که آن نیز توسط نقدریان و قبادی نوشته شده می پردازیم در این مقاله از نمودارهای کنترل فازی چند متغیره برای پایش پروفایل فاز ۱ در پایش عرض از مبدأ و شیب استفاده کردند. نتایج تحقیقی که براساس این نمودار بر روی داده های گردشگری انجام شد که عملکرد روش های جمع تجمعی چندمتغیره فازی $F-MCUSUM^{16}$ و میانگین متحرک موزون نمایی چندمتغیره فازی $F-MEWMA^{17}$ هنگامی که تغییر در نمونه های پایانی اتفاق می افتد نشان می دهد. اما روش $F-T_2$ دارای حساسیت یکنواخت نسبت به تغییرات در تمام محدوده تغییر است. عملکرد سه روش در تشخیص بزرگ در همه ی نقاط در نمونه ها کاملاً یکسان است. همچنین روش $I-MR$ نسبت به تغییرات کوچک حساس نیست از اینرو سه روش استفاده شده در این مقاله دارای قدرت و توانایی بیشتری نسبت به روش $I-MR$ در شناسایی تغییرات در اندازه های مختلف در تمام نقاط نمونه ها دارد. شیف در شیب نتایج تحقیق نشان می دهد روش $F-MEWMA$ کمی بهتر از روش $F-MCUSUM$ در برخی از تغییرات کوچک در نمونه ها است. با اینحال روش $F-MCUSUM$ بطور یکنواخت فازی بهتر از روش $F-MEWMA$ در تمام محدوده ی تغییرات است. نمودار $F-T_2$ عملکرد ضعیفی دارد. در پایان روش MR نسبت به سه روش دیگر عملکردی ضعیف تر را داراست. مطالعه ای در رابطه با رضایت مشتری برای نشان دادن کاربرد روش و بیان تحلیل حساسیت پارامترها برای ایجاد یک پروفایل فازی ارائه شده است.

(۳) توان روش پایش تک متغیره را برحسب احتمال هشدار محاسبه کرده و استفاده از رویکرد خود برای شناسایی تغییرات متوسط و بزرگ در پروفایل فرآیند را توصیه کردند. آنها یک چارچوب چند متغیره برای پایش پروفایل کیفیت محصول در فاز ۱ در یک محیط فازی ایجاد کردند که در آن ویژگی کیفیت پاسخ فازی است. سه روش کنترل چند متغیره برای نظارت بر فرایندها توسعه داده شد که یک پروفایل خطی فازی را نشان می دهد. عملکرد روش پیشنهادی بر مبنای احتمال هشدار در حالت های مختلف خارج از کنترل مورد بررسی قرار گرفت. مطالعات شبیه سازی نشان داد که روش های $F-MCUSUM$ و $F-MEWMA$ عملکرد خوبی در تشخیص تغییرات در پارامترهای پروفایل، به ویژه تغییرات کوچک و متوسط دارند و به طور یکنواخت نسبت به روش های $F-T_2$ تحت تغییرات گام مختلف در مکان های مختلف عمل می کنند. علاوه بر این، نشان داده شد که نه تنها مقدار تغییر بلکه مکان تغییر بر عملکرد نمودار کنترل $F-T_2$ تاثیر می گذارد. با این وجود، در مقایسه با روش تک متغیری $I-MR$ ، $F-T_2$ عملکرد بهتری در تشخیص تغییرات اندازه بزرگ دارند. در انتها یک مطالعه موردی در صنعت گردشگری برای نشان دادن قابلیت کاربرد روش پیشنهادی ارائه گردید. در سومین مقاله ی مورد بررسی که توسط (۷) به تحریر درآمده است روشی برای پایش پروفایل های خطی ساده با پاسخ غیر دقیق و مبهم ارائه شده است. دو روش جدید براساس آمار $F-T_2$ و $F-MEWMA$ برای پایش فاز ۲ پروفایل های خطی ساده فازی پیشنهاد شده است. در این مقاله مثالی از صنعت سرمایه و کاشی سازی آورده شده است. مطالعه شبیه سازی برای ارزیابی عملکرد روش های پیشنهادی معیار مقایسه (ARL) بیان می کند که روش های پیشنهادی در تشخیص تغییرات سائز مختلف در پروفایل های فرآیند بسیار موثر هستند. در این مقاله، روش جدیدی براساس آماره $F-T_2$ و $F-EWMA$ برای پایش

¹⁵Individuals- Moving Range

¹⁶Fuzzy- Multivariate Cumulative Sum

¹⁷Fuzzy- Multivariate Exponentially Weighted Moving Average

فاز ۲ پروفایل‌های خطی ساده فازی پیشنهاد شده‌است. عملکرد روش‌های ارائه شده ارزیابی و با استفاده از معیار خارج از کنترل (ARL) مقایسه می‌شوند. نتایج با استفاده از مجموعه داده‌های شبیه‌سازی شده نشان می‌دهد که عملکرد روش مبتنی بر آمار $F - EWMA$ نسبت به $F - T2$ برای تشخیص تغییرات کوچک و متوسط بهتر است. حال آنکه $F - T2$ در تشخیص تغییرات بزرگ بهتر عمل می‌کند. چهارمین مقاله ی مورد بررسی توسط (۱۱) نوشته شده‌است. در این مقاله، یک نمودار کنترل میانگین متحرک موزون نمایی فازی $F - EWMA$ جدید برای پایش پروفایل فاز ۲ فازی پیشنهاد شده‌است. در این راستا، نمودارهای کنترل $EWMA$ را به نوع فازی معادل خود بسط می‌دهیم. روش پیشنهادی قادر به شناسایی تغییرات کوچک در فرآیند تحت شرایط پایش پروفایل و توضیح پارامترهای مبهم، سخت و نادقیق است. تعیین منبع تغییرات، این روش ما را با امکان تشخیص علل انتقال فرآیند از حالت پایدار، حذف این علل و بازیابی حالت پایدار فراهم می‌کند. با توجه به این واقعیت که در بسیاری از مسائل دنیای واقعی، داده‌ها مبهم و نادقیق هستند، با علم به اینکه به‌کارگیری روش‌های کلاسیک کنترل فرآیند برای داده‌های نادقیق منطقی نیست، نظریه مجموعه فازی ابزاری کارآمد برای بررسی این کاستی است. مقدم و همکاران نشان دادند که این روش در کشف علل قابل تعیین در پروفایل بسیار موثر بوده‌است. آنها روش‌های رتبه‌بندی فازی را برای تعیین اینکه آیا پروفایل فازی فرآیند تحت کنترل است یا خیر استفاده کردند. روش رتبه‌بندی قادر به تشخیص تغییرات کوچکی در فرآیند است در شرایط غیرفازی و ثابت از یک نمونه بر نمونه‌ی باکلی برای محاسبه‌ی پروفایل هر نمونه‌ی تصادفی استفاده شده است. بر اساس این روش با داشتن بازه‌ی α برش می‌توانیم معادله‌ی خط فازی را به‌دست آوریم. می‌خواهیم بدانیم آیا پروفایل فرآیند در تحلیل تحت کنترل است یا خیر؟ چندین روش برای پایش پروفایل فاز ۲ وجود دارد. در این مقاله از آماره‌ی $F - EWMA$ ، که در واقع عدد مثلث فازی است، استفاده شده است. پارامتر موردنظر آماره‌ی $F - EWMA$ به گونه‌ای انتخاب می‌شود که نمودار کنترل ARL تحت کنترل خاص را حاصل کند برای محاسبه و رسم نمودار $F - EWMA$ از حساب جبری α برش‌ها و استفاده کرده‌اند. h حدکنترل نمودار یک مثلث فازی است اگر $F - EWMA > h$ و یا $F - EWMA < h$ باشد. نمودار کنترل از کنترل خارج می‌شود. در رگرسیون ما سه پارامتر داریم عرض از مبدأ، شیب، و واریانس که برای هر یک فاصله‌ی اطمینانی وجود دارد، ما از طریق فاصله‌ی اطمینان قادر خواهیم بود آنها را فازی کنیم. در این مقاله از روش رتبه‌بندی برای مقایسه اعداد فازی مثلثی استفاده کرده‌است. در واقع برای محاسبه‌ی $F - EWMA$ و اعداد فازی مثلثی باید دسته‌بندی شوند. از آنجا که در دنیای واقعی بسیاری از داده‌ها مبهم و نادقیق هستند.

جدول ۱: نمودار کنترل پروفایلیهای فازی

روش تشخیص تحت کنترل بودن	روش نمونه گیری	فاز	بر اساس آلفا برش	معیار مقایسه	نمودار کنترل فازی	نویسنده/نویسندگان (سال)	مثال کاربردی	نام مقاله
غیر فازی سازی و رتبه بندی فازی	SRS	۱	+	-	I-MR	تقدیریان و قبادی (۲۰۱۲)	صنعت گرومتری	profiles quality fuzzy of monitoring phase-I to approach univariate a Developing
غیر فازی سازی	bootstrap	۱	-	ARL	F-MEWMA F-MCUSUM T _Y - F	چندمنظیره (۲۰۱۴)	صنعت گرومتری	profiles quality fuzzy monitor to approach multivariate a Developing
رتبه بندی فازی	SRS	۲	+	ARL	F-EWMA F-T _Y	مقدم و همکاران (۲۰۱۵)	سرامیک و کاشی	PROFILES LINEAR FUZZY II PHASE monitor to methods new Developing
رتبه بندی فازی	SRS	۲	+	ARL	F-EWMA	مقدم و همکاران (۲۰۱۶)	دادهای شبیه سازی شده	profiles fuzzy II phase monitoring for charts control EWMA fuzzy New

[۱] بدخشان، ف. (۱۳۹۶). پایش پروفایل‌های خطی ساده هنگام نابرابری واریانس، دانشگاه فردوسی، مشهد.

[۲] مظفری پیمان، م. (۱۳۹۲). نمودارهای کنترل در محیط فازی، دانشگاه فردوسی، مشهد.

- [3] Noghondarian K and Ghobadi S. (2014), Developing a multivariate approach to monitor fuzzy quality profiles *Quality & Quantity*, 2 (2014), 817–836.
- [4] Lowry C.A and Montgomery D.C.(1995) A review of multivariate control charts, *IIE Trans* 26, 800–810 (1995).
- [5] Midi H. (2011), Robust multivariate control charts to detect small shifts in mean, *Math. Probl. Eng* 923463 (2011).
- [6] Noghondarian K and Ghobadi S.(2012) Developing a univariate approach to phase-I monitoring of fuzzy quality profiles, *Industrial Engineering Computations* 26, 800–810 (1995).
- [7] Moghadam G, Raissi Ardali G, & Amirzadeh V. (2015) Developing new methods to monitor phase II fuzzy linear profiles, *Iranian Journal of fuzzy systems* 4, 59-77 (2015).
- [18] Zadeh L.A. (1965) Fuzzy sets, *Information and contro* 3, 338-353 (1965).
- [9] Mahmoud M. A and Woodall W. H.(2004), Phase I analysis of linear profiles with calibration applications, *Technometrics* 46, 380-391 (2004).
- [10] Mestek O, Pavlik J, & Suchánek M. (1994), Multivariate control charts: control charts for calibration curves, *Frese-nius'J. Anal. Chem* 350, 344–351 (1994).
- [11] Moghadam G, Raissi Ardali G, & Amirzadeh V. (2015) New F-EWMA control charts for monitoring phase II fuzzy profiles, *Decision Science Letters* 5, 119–128 (2016).
- [12] Bradshaw C.W. (1983) A fuzzy set theoretic interpretation of economic control limits, *European Journal of Operational Research* 4, 403–408 (1983).
- [13] Sabegh , Zavvar M.H, Mirzazadeh A,& Salehian S, Weber G. (2014) A Literature Review on the Fuzzy Control Chart; Classifications & Analysis, *International Journal of Supply and Operations Management* 2, 167-189 (2014).

نمودار RA-CUSUM بر اساس داده‌های فازی

افسانه رضایی^۱، دکتر بهرام صادق‌پور گیلده^۲، دکتر غلامرضا محتشمی برزادران^۳

^{۱،۲،۳} گروه آمار، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد

چکیده

نمودار جمعی تجمعی ریسک تعدیل‌شده (Risk adjusted CUSUM) برای پایش فرآیندهای جراحی در بیمارستان‌ها و مراکز درمانی ارائه شد و به دلیل در نظر گرفتن ریسک عمل جراحی بیماران مورد توجه قرار گرفت. در این مطالعه فرض شده است که ریسک یک عدد مبهم و فازی باشد و بر این اساس نمودار RA-CUSUM بر اساس داده‌های فازی معرفی می‌شود. سپس نتایج بر روی داده‌های حاصل از عمل جراحی بیماران قلبی بررسی می‌شود.

کلمات کلیدی: نسبت بخت‌ها، ریسک بیماران، رگرسیون لجستیک فازی، RA-CUSUM.

۱ مقدمه

نمودارهای CUSUM نخستین بار توسط (۸) برای بررسی فرآیندهای صنعتی معرفی شد. اگرچه این نمودار نسبت به تغییرات کوچک حساس است اما در حوزه بهداشت و درمان و در کنترل فرآیندهای جراحی؛ بیماران به دلیل داشتن ویژگی‌های متفاوت از قبیل سن، جنسیت، سابقه بیماری و ... جامعه ناهمگونی را تشکیل می‌دهند و باید ریسک‌های مرتبط با ویژگی افراد را در نظر گرفت بنابراین نمودار CUSUM نمی‌تواند مناسب باشد. از این رو نمودارهای ریسک تعدیل‌شده در این زمینه ارائه شدند. (۵؛ ۶) و (۱۲) مدلی بر اساس تفاوت بین مرگ و میر واقعی و مورد انتظار به صورت تجمعی ارائه دادند. اما در مدل لاوگرو مشخص نمی‌شود که چه مقدار اختلاف نگران‌کننده است و در مدل پولونیکی حدود کنترل در تغییر است و بروزرسانی مکرر باعث عدم تشخیص تغییرات تدریجی مدل می‌شود. پس از آن مدل CUSUM ریسک تعدیل‌شده (RA-CUSUM) توسط (۷) معرفی شد و (۴) یک مدل RA-CUSUM کلی معرفی کرد به طوری که مدل استینر حالت خاصی از آن است. برای مشخص کردن ریسک بیمار در این مدل‌ها در زمانی که عملکرد جراحان بر روی بیماران با عمل قلبی مدنظر باشد از نمره پارسوننت ((۱۰))، استفاده می‌شود. با استفاده از آن و مدل رگرسیون لجستیک مناسب ریسک قبل از عمل جراحی قلبی بیماران پیش‌بینی می‌شود. با توجه به اینکه ریسک یک مقدار مبهم و نادقیق است بهتر است یک عدد فازی در نظر گرفته شود. در این حالت لازم است از رگرسیون فازی ((۸)) جهت تعیین ریسک قبل از عمل بیماران،

^۱ افسانه رضایی فر : afsaneh.rezaeefar@gmail.com

استفاده شود.

در این مقاله نمودار RA-CUSUM بر اساس داده‌های فازی را معرفی می‌کنیم. این نمودار می‌تواند به صورت دقیق‌تری فرآیند جراحی را در سه سطح α -برش بالا، α -برش پایین و زمانی که $\alpha = 1$ باشد بررسی کند و در زمانی که داده‌ها فازی باشند مناسب است.

در بخش ۲ خلاصه‌ای از مفاهیم فازی و رگرسیون لجستیک فازی بیان می‌شود. در بخش ۳ مدل RA-CUSUM بر اساس داده‌های فازی LR تعریف و بررسی می‌شود و در بخش ۴ یک مثال کاربردی جهت بررسی مدل معرفی شده، ارائه می‌شود.

۲ رگرسیون لجستیک با داده‌های فازی LR

در ابتدا به صورت مختصر اعداد فازی و برخی از ویژگی‌های آن بیان و سپس رگرسیون لجستیک فازی معرفی می‌شود.

به منظور تجزیه و تحلیل مجموعه داده‌هایی که دارای ویژگی مبهم و نادقیق هستند، مفاهیم نظریه مجموعه فازی به کار گرفته می‌شود. در واقع اگر $A \subseteq X$ یک مجموعه فازی باشد؛ آن‌گاه یک زیرمجموعه فازی از آن توسط یک تابع $\mu_A(x) : X \rightarrow [0, 1]$ به نام تابع عضویت مشخص می‌شود که برای هر $x \in X$ ، میزان عضویت x در مجموعه فازی A را نشان می‌دهد. ((۱))

تعریف ۱.۰۲. مجموعه‌ای از X که درجه عضویت آن در مجموعه فازی A دست کم به بزرگی α ($0 \leq \alpha \leq 1$) باشد؛ α -برش A گوئیم و با A_α نشان می‌دهیم؛ یعنی

$$A_\alpha = \{x \in X \mid \mu_A(x) \geq \alpha\}.$$

مجموعه فازی $A \subseteq R$ را یک عدد فازی گوئیم هرگاه μ_A نرمال و تک‌نمایی باشد و α -برش‌های آن به ازای هر $\alpha \in [0, 1]$ به صورت بازه‌های بسته باشد.

تعریف ۲.۰۲. فرض کنید $f : X \rightarrow Y$ یک تابع و A یک مجموعه فازی از X باشد. در این صورت $B = f(A)$ به صورت یک مجموعه فازی از Y با تابع عضویت زیر تعریف می‌شود.

$$\mu_B(y) = \begin{cases} \sup \mu_A(x)_{y=f(x); x \in X}, & f^{-1}(\{y\}) \neq \emptyset, \\ 0, & f^{-1}(\{y\}) = \emptyset. \end{cases}$$

که در آن $f^{-1}(\{y\})$ نگاشت معکوس f است. این تعریف اصل توسعه نام دارد.

تعریف ۳.۰۲. اگر ساختار تابع عضویت عدد فازی A به صورت زیر باشد:

$$\mu_A(x) = \begin{cases} L\left(\frac{m-x}{\alpha}\right), & x \leq m, \\ R\left(\frac{x-m}{\beta}\right), & x > m. \end{cases}$$

که در آن L و R توابع غیرصعودی از $[0, 1] \rightarrow R^+$ هستند و $R(0) = L(0) = 1$ ؛ آن‌گاه A را یک عدد فازی LR با نماد $(m, \alpha, \beta)_{LR}$ می‌نامیم؛ که در آن m مقدار نما و اعداد مثبت α و β به ترتیب پهناهای چپ و راست است.

در صورتی که $L(x) = R(x) = \max\{0, 1 - |x|\}$ ؛ آن‌گاه A یک عدد فازی مثلثی به صورت $(m, \alpha)_T$ است.

قضیه ۴.۰۲. اگر $M = (m, \alpha, \beta)_{LR}$ ، $N = (n, \delta, \gamma)_{LR}$ و $\lambda \in \mathbb{R}$ باشد، آن‌گاه

$$\lambda \otimes (m, \alpha, \beta)_{LR} = \begin{cases} (\lambda m, \lambda \alpha, \lambda \beta)_{LR}, & \lambda > 0, \\ (\lambda m, -\lambda \beta, -\lambda \alpha)_{LR}, & \lambda < 0. \end{cases} \quad (\text{الف})$$

$$M \oplus N = (m + n, \alpha + \delta, \beta + \gamma)_{LR} \quad (\text{ب})$$

$$M \ominus N = (m - n, \alpha + \gamma, \beta + \delta)_{LR} \quad (\text{ج})$$

(د) در صورتی که M و N دو عدد فازی مثبت باشند؛

$$M \oslash N \simeq \left(\frac{m}{n}, \frac{\gamma m + \alpha n}{n^2}, \frac{\delta m + \beta n}{n^2} \right)_{LR}$$

که در آن علامت‌های $\otimes, \odot, \ominus$ و $\oplus$ به ترتیب عملگرهای ضرب، تقسیم، تفریق و جمع برای اعداد فازی هستند. ذکر این نکته لازم است که حاصل تقسیم دو عدد LR در حالت کلی LR نیست، اما روابط تقریبی فوق برای این عملگر پیشنهاد شده است.

در ادامه رگرسیون لجستیک فازی معرفی می‌شود.

در مدل رگرسیون لجستیک فازی با روش کمترین مربعات فرض بر این است که پاسخ مشاهده‌شده فازی باشد و پاسخ برآورده‌شده نیز فازی خواهد بود. این مدل به صورت زیر تعریف می‌شود

$$\tilde{W}_i = \ln \left(\frac{y_i}{1 - y_i} \right) = \tilde{b}_0 + \tilde{b}_1 x_{i1} + \dots + \tilde{b}_m x_{im}. \quad (1)$$

که در آن $\tilde{b}_0, \tilde{b}_1, \dots, \tilde{b}_m$ پارامترهای مدل و $X_i = (x_{i1}, \dots, x_{im})$ بردار مشاهدات دقیق متغیر مستقل برای i امین مورد می‌باشد. در این صورت پاسخ برآورد شده فازی به صورت زیر است.

$$\tilde{W}_i = (f_i(a), f_i(s))_T, \quad (2)$$

که در آن

$$f_i(a) = a_0 + a_1 x_{i1} + \dots + a_m x_{im},$$

و

$$f_i(s) = s_0 + s_1 x_{i1} + \dots + s_m x_{im}.$$

با استفاده از روش کمترین مربعات مقادیر $a_0, a_1, \dots, a_m$ و $s_0, s_1, \dots, s_m$ برآورد می‌شوند.

در انتها با استفاده از اصل توسیع، y_i در مدل (۱) به صورت زیر به دست می‌آیند.

$$\tilde{y}_i = \left(\frac{e^{f_i(a)}}{1 + e^{f_i(a)}}, \frac{e^{f_i(a)}}{1 + e^{f_i(a)}} - \frac{e^{f_i(a)-f_i(s)}}{1 + e^{f_i(a)-f_i(s)}}, \frac{e^{f_i(a)+f_i(s)}}{1 + e^{f_i(a)+f_i(s)}} - \frac{e^{f_i(a)}}{1 + e^{f_i(a)}} \right)_{LR} \quad (3)$$

۳ نمودار RA-CUSUM بر اساس داده‌های فازی

نمودار RA-CUSUM توسط استینر و بر اساس نسبت بخت معرفی شد، به طوری که

$$\begin{cases} H_0 = \frac{d_0(y_n)}{1 - d_0(y_n)} = Q_0 \frac{y_n}{1 - y_n}, \\ H_A = \frac{d_A(y_n)}{1 - d_A(y_n)} = Q_A \frac{y_n}{1 - y_n}. \end{cases} \quad (4)$$

که در آن $d_{\circ}$ و d_A به ترتیب احتمال مردن بیمار تحت فرضیه $H_{\circ}$ و H_A و y_n ریسک بیمار n ام است. ضریب $Q_{\circ}$ معمولاً یک و Q_A برحسب اینکه بخواهیم پسرفت ($Q_A > 1$) و پیشرفت ($Q_A < 1$) عملکرد جراحان را تعیین کنیم مشخص می‌شود. با توجه به (۴) داریم

$$d_{\circ}(y_n) = \frac{Q_{\circ} y_n}{1 - y_n + Q_{\circ} y_n},$$

$$d_A(y_n) = \frac{Q_A y_n}{1 - y_n + Q_A y_n}.$$

و آماره جمعی تجمعی S_n به صورت زیر تعریف می‌شود.

$$S_n = \max\{\circ, S_{n-1} + W_n\}, \quad (5)$$

که در آن W_n بر اساس لگاریتم نسبت درستنمایی به صورت زیر به دست می‌آید.

$$W_n = \begin{cases} \log \frac{(1-y_n+Q_{\circ}y_n)Q_A}{(1-y_n+Q_Ay_n)Q_{\circ}}, & \text{اگر بیمار } n \text{ ام بمیرد.} \\ \log \frac{(1-y_n+Q_{\circ}y_n)}{(1-y_n+Q_Ay_n)}, & \text{اگر بیمار } n \text{ ام زنده بماند.} \end{cases}$$

و با رسم S_n در برابر n نمودار RA-CUSUM به دست می‌آید.

در صورتی که $S_n > h$ حد کنترل است؛ نمودار هشدار می‌دهد.

اگر فرض کنیم ریسک بیمار به صورت عدد فازی LR باشد. $(\tilde{y}_n = (m, l_n, r_n)_{LR})$ ؛ آن‌گاه فرضیه $H_{\circ}$ و H_A به صورت زیر تغییر می‌یابد.

$$\begin{cases} H_{\circ} = \frac{\tilde{d}_{\circ}(y_n)}{1 - \tilde{d}_{\circ}(y_n)} = Q_{\circ} \frac{\tilde{y}_n}{1 - \tilde{y}_n}, \\ H_n = \frac{\tilde{d}_A(y_n)}{1 - \tilde{d}_A(y_n)} = Q_A \frac{\tilde{y}_n}{1 - \tilde{y}_n}. \end{cases}$$

در این صورت بر اساس حساب اعداد فازی LR داریم:

$$\frac{\tilde{d}_{\circ}(y_n)}{1 - \tilde{d}_{\circ}(y_n)} = \left(\frac{Q_{\circ} m_n}{1 - m_n}, \frac{l_n}{(1 - m_n)^2}, \frac{r_n}{(1 - m_n)^2} \right)_{LR}.$$

در این حالت آماره زیر را خواهیم داشت:

$$\tilde{S}_n(s) = \max_{s=\max\{a,y\}} \{\circ(a), \tilde{S}_{n-1} \oplus \tilde{W}_n(y)\}, \quad s \in \mathbb{R}.$$

که در آن صفر یک عدد فازی LR با پهنای صفر است و

$$\tilde{W}_n = \begin{cases} \log \left\{ \frac{(1 \ominus \tilde{y}_n \oplus Q_{\circ} \odot \tilde{y}_n) \odot Q_A}{(1 - \tilde{y}_n \oplus Q_A \odot \tilde{y}_n) \odot Q_{\circ}} \right\}, & \text{اگر بیمار } n \text{ ام بمیرد,} \\ \log \left\{ \frac{1 \ominus \tilde{y}_n \oplus Q_{\circ} \odot \tilde{y}_n}{1 \ominus \tilde{y}_n \oplus Q_A \odot \tilde{y}_n} \right\}, & \text{اگر بیمار } n \text{ ام زنده بماند.} \end{cases}$$

در نتیجه $\tilde{W}_n$ یک عدد فازی است و α -برش‌های آن به صورت زیر تعریف می‌شود.

$$W_n^-(\alpha) = \begin{cases} \log \left\{ \frac{Q_A}{1 - m_n + Q_A m_n} (C_1(\alpha - 1) + 1) \right\}, & \text{اگر بیمار } n \text{ ام بمیرد,} \\ \log \left\{ \frac{1}{(1 - m_n + Q_A m_n)} (C_1(\alpha - 1) + 1) \right\}, & \text{اگر بیمار } n \text{ ام زنده بماند.} \end{cases}$$

$$W_n^+(\alpha) = \begin{cases} \log \left\{ \frac{Q_A}{(1 - m_n + Q_A m_n)} (1 + (1 - \alpha)C_2) \right\}, & \text{اگر بیمار } n \text{ ام بمیرد,} \\ \log \left\{ \frac{1}{(1 - m_n + Q_A m_n)} (1 + (1 - \alpha)C_2) \right\}, & \text{اگر بیمار } n \text{ ام زنده بماند.} \end{cases}$$

که در آن $\alpha \in [0, 1]$ و

$$C_1 = \frac{(l_n + Q_A r_n) + (r_n + l_n)(1 - m_n + Q_A m_n)}{(1 - m_n + Q_A m_n)},$$

و

$$C_2 = \frac{(r_n + Q_A l_n) + (r_n + l_n)(1 - m_n + Q_A m_n)}{(1 - m_n + Q_A m_n)}.$$

در نتیجه

$$S_n^-(\alpha) = \max\{0, S_{n-1}^-(\alpha) + W_n^-(\alpha)\}.$$

و

$$S_n^+(\alpha) = \max\{0, S_{n-1}^+(\alpha) + W_n^+(\alpha)\}.$$

که در آن $\alpha \in [0, 1]$ و حدود کنترل برای هر کدام از α -برش‌های بالا و پایین مشخص می‌گردد؛ به طوری که اگر $S_n^-(\alpha) > h^-(\alpha)$ یا $S_n^+(\alpha) > h^+(\alpha)$ هشدار رخ می‌دهد.

برای به دست آوردن حدود کنترل لازم است توزیع هسته ریسک (m_n) و پهنای چپ و راست آن (r_n, l_n) را بدانیم. پس از مشخص شدن توزیع آن‌ها و خطای نوع اول مدل (و متوسط طول اجرا)، حدود کنترل بالا و پایین با استفاده از شبیه‌سازی و تولید تصادفی از توزیع‌های مشخص شده به دست می‌آید.

۴ مثال

با استفاده از داده‌های مقاله (۷) که نتایج عمل قلبی ۶۹۹۴ بیمار در سال‌های ۱۹۹۲-۱۹۹۸ میلادی در یک مرکز جراحی در انگلستان است؛ مدل معرفی شده بررسی می‌گردد. در این داده‌ها تعداد افراد فوت شده تا ۳۰ روز پس از عمل جراحی ۴۶۱ نفر با نرخ مرگ و میر ۶/۶٪ بوده است و نمره پارسوننت بیماران بر اساس مشخصات افراد از قبیل جنسیت، سن، سابقه بیماری و ... توسط متخصصین مربوطه مشخص شده است و می‌تواند متغیر توضیحی مناسبی جهت پیش‌بینی ریسک عمل بیماران باشد؛ برای پیدا کردن پارامترهای مدل در فاز ۱ دو سال اول داده‌ها (۱۹۹۲-۱۹۹۳) استفاده گردید و شروع مدل در فاز ۲ از ابتدا سال ۱۹۹۴ در نظر گرفته شده است. مدل رگرسیون لجستیک فازی (۱) و (۲) به صورت زیر به دست آمد.

$$f_i(a) = -3/528 + 0/0554X_i,$$

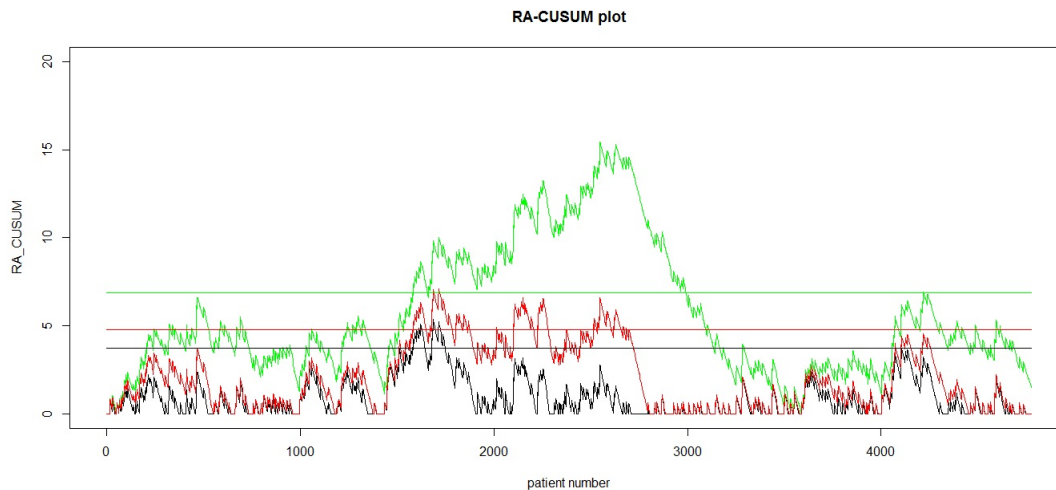
$$f_i(s) = 0/1834 + 0/000014X_i.$$

که در آن X_i نمره پارسوننت بیماران است.

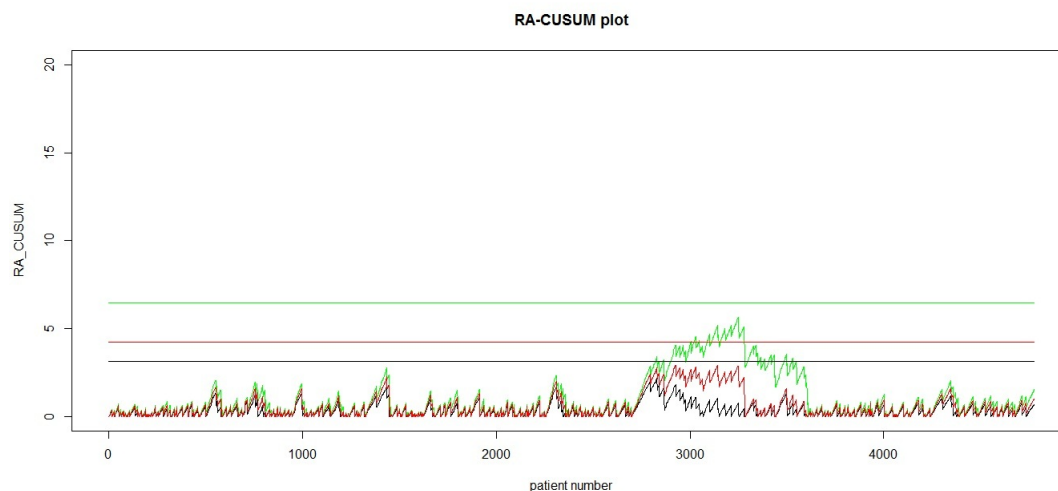
در اینجا عملکرد جراحان را با استفاده از مدل RA-CUSUM بر اساس داده‌های فازی با فرض $Q_0 = 1$ و در دو حالت $Q_A = 2$ (پسرفت عملکرد جراحان) و $Q_A = 0/5$ (پیشرفت عملکرد جراحان) بررسی می‌کنیم. نمودارهای (۱) و (۲) نمودار RA-CUSUM فازی را با فرض α -برش ۰/۸۵ در دو حالت فوق نشان می‌دهد.

با فرض $ARL_0 = 9600$ حدود کنترل با α -برش فوق به صورت $h^+(0/85) = 6/5$ و $h^-(0/85) = 3/5$ به دست آمد. در صورتی که α -برش برابر یک باشد $h = 4/5$ که همان مقداری است که استینر به دست آورد.

نمودارهای (۱) و (۲) عملکرد کل جراحان قلب مرکز را نشان می‌دهد. همان‌طور که در نمودار (۱) مشاهده می‌شود اواسط سال ۱۹۹۵ لغایت اواخر سال ۱۹۹۶ افت عملکرد پزشکان وجود داشته است و نمودار (۲) نشان می‌دهد که عملکرد پزشکان از اوایل سال ۱۹۹۷ لغایت اواسط سال ۱۹۹۷ پیشرفت داشته است.



شکل ۱:



شکل ۲:

۵ بحث و نتیجه‌گیری

در این مطالعه مدل RA-CUSUM زمانی که داده‌ها فازی باشند معرفی شد. از آنجایی که ریسک عمل بیماران عددی مبهم است و نمی‌تواند به صورت دقیق مشخص گردد؛ بهتر است آن را به صورت یک عدد فازی در نظر گرفته و با استفاده از مدل رگرسیون لجستیک فازی مناسب تعیین گردد. در این حالت نمودار RA-CUSUM متناظر با آن نیز تغییراتی پیدا کرده و آماره جمعی تجمعی و حدود کنترل آن به صورت یک عدد فازی تبدیل می‌شود و نمودار RA-CUSUM بر اساس داده‌های فازی را تشکیل می‌دهد. بنابراین زمانی که ریسک عمل بیماران یک عدد فازی باشد؛ این مدل می‌تواند جهت پایش عملکرد جراحان مناسب باشد.

مراجع

[۱] طاهری، م و ماشین چی، م. (۱۳۸۷). مقدمه‌ای بر احتمال و آمار فازی، چاپ اول، دانشگاه شهید باهنر، کرمان.

- [4] Gan F., Lin L. & Chok K. (2012). Risk-adjusted CUSUM Charting Procedures. *Frontiers in Statistical Quality Control* 10, 207-225.
- [5] Lovegrove, J., Valencia, O. Treasure, T., Sherlaw-Jahson, C., & Gallivan, S. (1997). Monitoring the result of Cardiac surgery by variable-life-adjusted display. *The Lancet*, 350, 1128-1130.
- [6] Lovegrove, J., Sherlaw-Jahson, C., Valencia, O., & Gallivan, S. (1999). Monitoring the Performance of Cardiac Surgeons. *Journal of the Operational Research Society*, 50, 685-689.
- [8] Page, E. S. (1954). Continuous inspection Schemes. *Biometrika*, 41, 100-115.
- [10] Parsonnet, V., Dean, D. & Bernstein, A. D. (1989). A method of uniform stratification of risks for evaluating the result surgery in acquired adult heart disease. *Circulation*, 779 (Supp11), 1-12.
- [12] Poloniecki, J., Valencia, O., & Littlejohns, P. (1998). Cumulative risk-adjusted mortality chart for detecting changes in death rate: Observational study of heart surgery. *British Medical Journal*, 316, 1697-1700.
- [8] Pourahmad S, Ayatollahi S. M. T., Taheri S. M. & Agahi Z. H. (2011). Fuzzy logistic regression based on the least squares approach with application in clinical studies. *Computers and Mathematics with applications*. 2011, 62:3353-3365.
- [7] Steiner, S. H., Cook, R. J., Farewell, V. T. & Treasure, T. (2000). Monitoring surgical performances using risk-adjusted cumulative sum charts. *Biostatistics*, 1, 441-452.

ضریب همبستگی در شرایط نادقیق

رضا زارعی^۱، طاهره کیانپور^۲

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

^۲ دانش‌آموخته کارشناسی ارشد آمار ریاضی، دانشکده علوم ریاضی، دانشگاه گیلان

چکیده

همبستگی یکی از مهم‌ترین شاخص‌هایی است که به‌طور گسترده در زمینه‌های مختلف به‌کار می‌رود و اندازه‌ای با اهمیت در تحلیل داده‌هاست. از آنجایی که بسیاری از داده‌های واقعی ممکن است نادقیق باشند، مفهوم همبستگی در محیط نادقیق توسعه داده شد. در این مقاله، ابتدا بر اساس تعریف کلاسیک ضریب همبستگی و اصل توسعه داده یک معیار برای محاسبه ضریب همبستگی بین اعداد فازی را مطالعه می‌کنیم. رویکرد برنامه‌ریزی ریاضی به‌منظور محاسبه درجه عضویت اندازه‌های فازی بکار گرفته شده است. سپس، روشی ساده برای تعیین و محاسبه دقیق تابع عضویت با استفاده از T -نرم دراستیک و اعمال جبری بر اساس آن مورد بررسی قرار گرفته است که به برنامه‌ریزی ریاضی بستگی ندارد. در پایان برخی روابط برای اندازه‌گیری ضرایب همبستگی بین مجموعه‌های فازی شهودی و مجموعه‌های فازی مردد مطالعه شده است.

کلمات کلیدی: اصل توسعه، برنامه‌ریزی ریاضی، ضریب همبستگی، عدد فازی، T -نرم دراستیک.

۱ مقدمه

بررسی ساختار همبستگی و ارتباط میان متغیرها یکی از مسایل مورد توجه در آمار است. در بیشتر تحلیل‌های آماری یافتن و تعیین نوع رابطه بین متغیرهای تصادفی و مشاهدات حاصل از آن‌ها همواره مورد توجه محققین بوده است. تاکنون اندازه‌های متنوعی تحت عنوان ضرایب همبستگی بدین منظور معرفی شده و روش‌های مختلفی برای بررسی این ارتباط وجود دارد که بسته به ویژگی داده‌های مورد بررسی و هدف تحقیق، می‌توان از هر کدام از آن‌ها استفاده نمود. بر حسب این‌که متغیرها کمی یا کیفی باشند، ضریب همبستگی شاخصی است که میزان رابطه بین متغیرها را نشان می‌دهد. از مهمترین ضرایب همبستگی که تاکنون معرفی شده‌اند می‌توان به ضریب همبستگی پیرسن، ضریب همبستگی رتبه‌ای اسپیرمن، C ضریب همبستگی توافق پیرسون یا ضریب توافق، ضریب همبستگی کرامر، ضریب همبستگی گاما، ضریب همبستگی فی اشاره کرد. یک ملاک مناسب برای تعیین همبستگی دو متغیر کمی ضریب همبستگی پیرسن می‌باشد که توسط پیرسون معرفی شد که از آن زمان تا کنون به

^۲ رضا زارعی، rezazarei.r@gmail.com, r-zarei@guilan.ac.ir

منظور تعیین میزان رابطه، نوع و جهت رابطه‌ی بین دو متغیر فاصله‌ای یا نسبی و یا یک متغیر فاصله‌ای و یک متغیر نسبی به کار برده می‌شود. با وجود کاربرد فراوان ضرایب همبستگی اشاره شده در بسیاری از مسایل واقعی به سبب شرایط حاکم بر مسئله خطای ماشین و یا خطای انسانی امکان اندازه‌گیری دقیق مشاهدات در مسئله تحت بررسی وجود ندارد. در این شرایط داده‌ها به صورت نادقیق گردآوری و ثبت شده، بنابراین استفاده از روش‌های دقیق برای تجزیه و تحلیل این گونه مسائل با خطاهای جبران‌ناپذیری همراه خواهد بود. از این‌رو نیازمندیم تا روش‌های جایگزین به منظور تجزیه و تحلیل همبستگی بین دو مجموعه از داده‌ها یا دو متغیر تصادفی در شرایطی که داده‌های گردآوری شده، نادقیق هستند ارائه دهیم. در این مقاله با استفاده از ابزارهای مفیدی که نظریه مجموعه‌های فازی برای صورت‌بندی مفاهیم نادقیق در اختیار ما قرار داده است به بررسی ضریب همبستگی تحت شرایط نادقیق خواهیم پرداخت.

۲ ضریب همبستگی بین اعداد فازی

با فرض این‌که $(X_1, Y_1), (X_2, Y_2), \dots, (X_n, Y_n)$ یک نمونه تصادفی از متغیرهای دو جامعه باشند، ضریب همبستگی پیرسن در این مشاهدات به صورت زیر تعریف می‌شود

$$\rho_{X,Y} = \frac{\sum_{i=1}^n (X_i - \bar{X})(Y_i - \bar{Y})}{\sqrt{\sum_{i=1}^n (X_i - \bar{X})^2 \sum_{i=1}^n (Y_i - \bar{Y})^2}}. \quad (1)$$

علاوه بر آن در صورتی‌که مقادیر مشاهده شده این نمونه تصادفی را با $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ نشان دهیم، برآورد ضریب همبستگی آن را با $r_{X,Y}$ نشان می‌دهیم که به صورت زیر محاسبه می‌شود

$$r_{X,Y} = \hat{\rho}_{X,Y} = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 \sum_{i=1}^n (y_i - \bar{y})^2}}. \quad (2)$$

با توجه به رابطه (۲)، به راحتی می‌توان نشان داد که $-1 \leq r_{X,Y} \leq +1$.

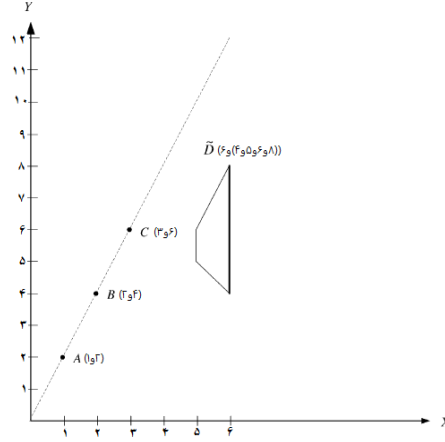
۱.۲ حالت خاص

در این بخش هدف محاسبه ضریب همبستگی بین زوج مشاهدات بر اساس یک نمونه تصادفی n تایی است که در بین آن‌ها یکی از داده‌ها دارای ابهام است و به صورت نادقیق مشاهده و ثبت شده است. به منظور سهولت در محاسبات فرض می‌کنیم که نمونه از چهار زوج مشاهده گردآوری و به صورت زیر ثبت شده‌اند (شکل ۱ را ببینید)

$$A = (1, 2), \quad B = (2, 4), \quad C = (3, 6), \quad D = (6, \tilde{Y}_D),$$

که در آن $\tilde{Y}_D = (4, 5, 6, 8)$ یک عدد فازی دوزنقه‌ای با تابع عضویت زیر می‌باشد

$$\mu_{\tilde{Y}_D}(y) = \begin{cases} y - 4 & 4 \leq y \leq 5, \\ 1 & 5 \leq y \leq 6, \\ \frac{8-y}{2} & 6 \leq y \leq 8. \end{cases}$$



شکل ۱: نمودار زوج مشاهدات A, B, C, D

برای تعیین مقدار ضریب همبستگی این مشاهدات نیازمند محاسبه $\bar{X}$ و $\bar{Y}$ می‌باشیم. واضح است که مقدار $\bar{X} = 3$ عددی دقیق است، اما با توجه به این‌که یکی از مشاهدات مربوط به Y_i ها نادقیق است، میانگین مشاهدات ثبت شده مربوط به Y_i ها به صورت زیر می‌باشد

$$\tilde{Y}_D = 3 + \frac{\tilde{Y}_D}{4}.$$

حال با استفاده از رابطه (۲) ضریب همبستگی به صورت زیر محاسبه می‌شود

$$\tilde{r}_{X,Y} = \frac{\sum_{i=A,B,C,D} (X_i - 3)(Y_i - 3 - \frac{\tilde{Y}_D}{4})}{\sqrt{\sum_{i=A,B,C,D} (X_i - 3)^2 \sum_{i=A,B,C,D} (Y_i - 3 - \frac{\tilde{Y}_D}{4})^2}}. \quad (3)$$

علی‌رغم این‌که عبارت به دست آمده در رابطه (۳) یک کمیت فازی است، لزومی ندارد که شکل تابع عضویت آن با تابع $\tilde{Y}_D$ (به عنوان تنها مشاهداتی که ابهام ایجاد کرده است) یکسان باشد، چرا که رابطه (۲) یک ترکیب غیر خطی از $\tilde{Y}_D$ است. برای به دست آوردن تابع عضویت $\tilde{r}_{X,Y}$ از عملگرهای محاسباتی که بر اساس $-\alpha$ برش اعداد فازی پایه‌گذاری شده‌اند، استفاده می‌کنیم. به این منظور ابتدا $-\alpha$ برش عدد فازی دوزنقه‌ای $\tilde{Y}_D$ ، $\tilde{r}_{X,Y}$ را به ترتیب به صورت زیر محاسبه می‌کنیم

$$\begin{aligned} (\tilde{Y}_D)_\alpha &= [(\tilde{Y}_D)_\alpha^L, (\tilde{Y}_D)_\alpha^U] \\ &= [\min_y \{y \in [4, 8] \mid \mu_{\tilde{Y}_D}(y) \geq \alpha\}, \max_y \{y \in [4, 8] \mid \mu_{\tilde{Y}_D}(y) \geq \alpha\}], \\ (\tilde{r}_{X,Y})_\alpha &= [(\tilde{r}_{X,Y})_\alpha^L, (\tilde{r}_{X,Y})_\alpha^U] \\ &= [\min_r \{r \in [-1, 1] \mid \mu_{\tilde{r}_{X,Y}}(r) \geq \alpha\}, \max_r \{r \in [-1, 1] \mid \mu_{\tilde{r}_{X,Y}}(r) \geq \alpha\}]. \end{aligned}$$

حال برای هر مقدار دلخواه $0 \leq \alpha \leq 1$ می‌توانیم $-\alpha$ برش‌های $(\tilde{Y}_D)_\alpha^L, (\tilde{Y}_D)_\alpha^U$ را که مقادیری دقیق هستند محاسبه کرد و با جایگذاری آن‌ها در رابطه (۳)، $-\alpha$ های برش $(\tilde{r}_{X,Y})_\alpha^L, (\tilde{r}_{X,Y})_\alpha^U$ را تعیین کرد.

برای تعیین تابع عضویت ضریب همبستگی $\tilde{r}$ ، ابتدا $-\alpha$ برش‌های عدد فازی دوزنقه‌ای $\tilde{Y}_D$ را به صورت زیر تعیین می‌کنیم

$$(\tilde{Y}_D)_\alpha = [4 + \alpha, 8 - 2\alpha],$$

با قرار دادن این مقادیر در رابطه (۳) $-\alpha$ های برش $\tilde{r}$ به صورت زیر خواهد آمد

$$(\tilde{r}_{X,Y})_\alpha^L = \frac{(4 + 3\alpha)}{\sqrt{10/5\alpha^2 + 112}}, \quad (\tilde{r}_{X,Y})_\alpha^U = \frac{(16 - 6\alpha)}{\sqrt{42\alpha^2 - 168\alpha + 280}}. \quad (4)$$

حال برای تعیین تابع عضویت $\tilde{r}$ از دو رابطه به دست آمده α - برش های آن استفاده می کنیم. با اندکی محاسبات جبری تابع عضویت $\tilde{r}$ به صورت زیر به دست می آید

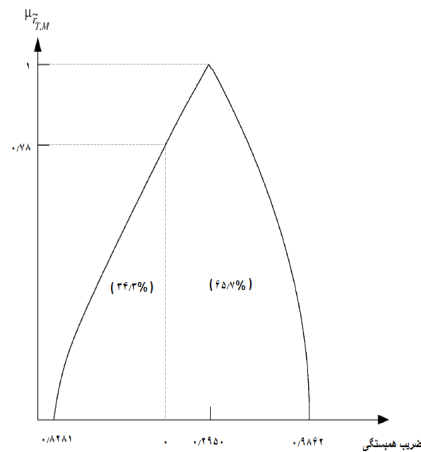
$$\mu_{\tilde{r}_{XY}}(r) = \begin{cases} \frac{-24+28\sqrt{6r^2-6r^4}}{18-21r^2} & \frac{1}{\sqrt{V}} \leq r \leq \sqrt{\frac{2}{D}}, \\ 1 & \sqrt{\frac{2}{D}} \leq r \leq \frac{1}{\sqrt{154}}, \\ \frac{48-42r^2-14\sqrt{6r^2-6r^4}}{18-21r^2} & \frac{1}{\sqrt{154}} \leq r \leq \frac{1}{\sqrt{V_0}}. \end{cases} \quad (5)$$

توجه داشته باشید که با توجه به اینکه $\tilde{Y}_D$ (تنها مشاهده دارای ابهام در مجموعه مشاهدات است) از نوع دوزنقه ای بوده، برای تعیین مقادیر ابتدایی، وسط و انتهایی به صورت زیر عمل می کنیم. دامنه سمت چپ و راست عدد فازی دوزنقه ای به ترتیب عبارتند از

$$[(\tilde{r}_{X,Y})_{\alpha=0}^L, (\tilde{r}_{X,Y})_{\alpha=1}^L] = [\frac{1}{\sqrt{V}}, \sqrt{\frac{2}{D}}],$$

$$[(\tilde{r}_{X,Y})_{\alpha=0}^U, (\tilde{r}_{X,Y})_{\alpha=1}^U] = [\frac{1}{\sqrt{154}}, \frac{1}{\sqrt{V_0}}].$$

همان طور که انتظار می رفت بر اساس این رابطه کمترین مقدار ضریب همبستگی برابر با $0/378$ و بیشترین مقدار ضریب همبستگی $0/9562$ است که دقیقاً با مقادیر به دست آمده از رابطه (۴) مطابقت دارد. با این وجود باید توجه نمود که هر چند تنها مشاهده نادقیق در مجموعه مشاهدات موجود، یک عدد فازی دوزنقه ای است، اما شکل حاصل از ضریب همبستگی فازی به دست آمده (شکل (۲) را ببینید) یک عدد فازی دوزنقه ای نیست.



شکل ۲: نمودار ضریب همبستگی بین مشاهدات A, B, C, D

۲.۲ حالت کلی

فرض کنید $(\tilde{X}_1, \tilde{Y}_1), (\tilde{X}_2, \tilde{Y}_2), \dots, (\tilde{X}_n, \tilde{Y}_n)$ نمونه ای از مشاهدات نادقیق جمع آوری شده دو متغیر تحت بررسی باشد که هدف محاسبه و تحلیل ضریب همبستگی بین آنها است. مشابه با رابطه (۲)، ضریب همبستگی این مشاهدات نادقیق به صورت زیر قابل بازنویسی است

$$\tilde{r}_{X,Y} = \frac{\sum_{i=1}^n (\tilde{X}_i - \frac{\sum_{i=1}^n \tilde{X}_i}{n})(\tilde{Y}_i - \frac{\sum_{i=1}^n \tilde{Y}_i}{n})}{\sqrt{\sum_{i=1}^n (\tilde{X}_i - \frac{\sum_{i=1}^n \tilde{X}_i}{n})^2 \sum_{i=1}^n (\tilde{Y}_i - \frac{\sum_{i=1}^n \tilde{Y}_i}{n})^2}}. \quad (6)$$

برخلاف بخش قبل که تنها یک داده نادقیق وجود داشت رابطه (۶) تابعی از اعداد فازی مختلف است، که محاسبه مقدار آن نیازمند انجام محاسبات جبری فراوان و پیچیده روی این اعداد است، که به سادگی نتیجه بخش نخواهد بود. در چنین شرایطی استفاده از اصل گسترش کار را بسیار ساده‌تر خواهد کرد. بر اساس این اصل تابع عضویت ضریب همبستگی این مشاهدات به صورت زیر بیان می‌شود.

$$\mu_{\tilde{r}_{X,Y}}(r) = \sup_{X,Y} \min\{\mu_{\tilde{X}_i}(x_i), \mu_{\tilde{Y}_i}(y_i), \forall i | r = r_{XY}\}, \quad (۷)$$

که در آن $r_{X,Y}$ در رابطه (۲) تعریف شده است.

با توجه به رابطه (۷)، α -برش‌های $\tilde{X}_i$ و $\tilde{Y}_i$ را (برای $\alpha \leq 1$) به صورت زیر در نظر بگیرید

$$\begin{aligned} (\tilde{X}_i)_\alpha &= [(\tilde{X}_i)_\alpha^L, (\tilde{X}_i)_\alpha^U] = [\min\{x_i \in X | \mu_{\tilde{X}_i}(x_i) \geq \alpha\}, \max\{x_i \in X | \mu_{\tilde{X}_i}(x_i) \geq \alpha\}], \\ (\tilde{Y}_i)_\alpha &= [(\tilde{Y}_i)_\alpha^L, (\tilde{Y}_i)_\alpha^U] = [\min\{y_i \in Y | \mu_{\tilde{Y}_i}(y_i) \geq \alpha\}, \max\{y_i \in Y | \mu_{\tilde{Y}_i}(y_i) \geq \alpha\}]. \end{aligned}$$

برای محاسبه تابع عضویت ضریب همبستگی فازی $\tilde{r}_{X,Y}$ بایستی کران‌های پایین و بالای α -برش این کمیت فازی را محاسبه کنیم که با توجه به توضیحات ارائه شده می‌توان این دو کران را در قالب دو برنامه‌ریزی غیرخطی به صورت زیر مدل بندی کرد

$$(\tilde{r}_{X,Y})_\alpha^L = \frac{\min \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 \sum_{i=1}^n (y_i - \bar{y})^2}}, \quad (۸)$$

$$s.t. \quad (\tilde{X}_i)_\alpha^L \leq x_i \leq (\tilde{X}_i)_\alpha^U, \quad \forall i, \quad (\tilde{Y}_i)_\alpha^L \leq y_i \leq (\tilde{Y}_i)_\alpha^U, \quad \forall i,$$

$$(\tilde{r}_{X,Y})_\alpha^U = \frac{\max \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum_{i=1}^n (x_i - \bar{x})^2 \sum_{i=1}^n (y_i - \bar{y})^2}}. \quad (۹)$$

$$s.t. \quad (\tilde{X}_i)_\alpha^L \leq x_i \leq (\tilde{X}_i)_\alpha^U, \quad \forall i, \quad (\tilde{Y}_i)_\alpha^L \leq y_i \leq (\tilde{Y}_i)_\alpha^U, \quad \forall i.$$

برای حل این برنامه‌ریزی‌های غیرخطی از روش گرادیان و نرم افزار متلب استفاده خواهیم کرد. برای دو مقدار α_1 و α_2 دلخواه که $0 < \alpha_2 < \alpha_1 \leq 1$ هستند، ناحیه جواب‌های شدنی تعریف شده توسط α_1 در برنامه‌های ارائه شده رابطه (۸) کوچکتر از ناحیه تعریف شده توسط α_2 خواهد بود. در نتیجه α -برش‌های $\tilde{r}_{X,Y}$ در سطح α_1 در α -برش‌های ایجاد شده در سطح α_2 قرار خواهد گرفت و به این دلیل تابع عضویت ضریب همبستگی $\tilde{r}_{X,Y}$ محدب خواهد شد. حال با داشتن شرط تحدب مورد بحث می‌توانیم پهنه‌های چپ و راست تابع عضویت $\tilde{r}_{X,Y}$ را با استفاده از α -برش‌های به دست آمده محاسبه کنیم و در نهایت به فرم تابع عضویت خواهیم رسید. نکته قابل ذکر این است که با توجه به تابع عضویت در نظر گرفته شده برای داده‌های نادقیق اولیه، ممکن است فرم پهنه‌های چپ و راست تابع عضویت، آنقدر پیچیده به دست آیند که امکان پیدا کردن خود تابع عضویت وجود نداشته باشد. در این شرایط از روش‌های عددی برای محاسبه میزان عضویت استفاده خواهیم کرد.

با در نظر گرفتن $L(r)$ و $R(r)$ به عنوان پهنه‌های چپ و راست تابع عضویت، می‌توان فرم نهایی این تابع را از رابطه زیر به دست آورد

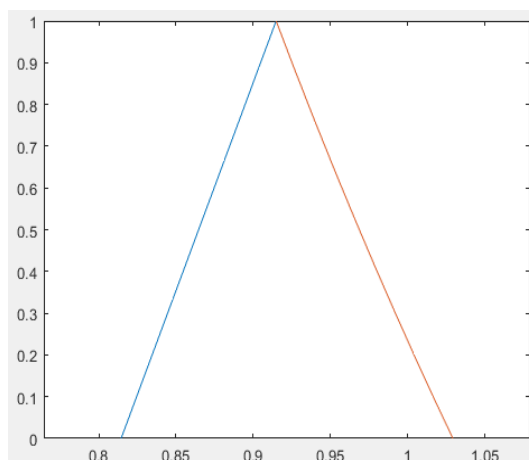
$$\mu_{\tilde{r}_{X,Y}}(r) = \begin{cases} L(r) & (\tilde{r}_{X,Y})_{\alpha=0}^L \leq r \leq (\tilde{r}_{X,Y})_{\alpha=1}^L, \\ 1 & (\tilde{r}_{X,Y})_{\alpha=1}^L \leq r \leq (\tilde{r}_{X,Y})_{\alpha=1}^U, \quad (1^0) \\ R(r) & (\tilde{r}_{X,Y})_{\alpha=1}^U \leq r \leq (\tilde{r}_{X,Y})_{\alpha=0}^U. \end{cases}$$

مثال ۱۰۲. رابطه بین شاخص‌های تکنولوژی و مدیریت همواره موضوعی مهم بوده و اندازه‌گیری آن مورد توجه محققین بسیاری قرار گرفته است. (۳) با استفاده از اطلاعات موجود راجع به این دو متغیر که غالباً به شکل متغیرهای زیانی هستند آن‌ها را در قالب اعداد فازی مدل‌بندی و ارائه نموده‌اند. بر این اساس، اطلاعات مربوط به سطوح مختلف تکنولوژی ($\tilde{T}_i$) و مدیریت ($\tilde{M}_i$) در یک نمونه ۱۵ تایی از شرکت‌های صنعتی کشور تایوان توسط این محققین گردآوری شد. بر اساس روش مورد مطالعه در این فصل و با استفاده از α -برش‌های موجود برای هر یک از سطوح مختلف برش، می‌توان یک برنامه‌ریزی غیرخطی تشکیل داد و با استفاده از نرم‌افزار متلب آن را حل کرد و جواب نهایی که همان α -برش مقدار ضریب همبستگی می‌باشد را به دست آورد. در پایان نتایج مربوط به هر یک از سطوح مختلف α در جدول (۱) ارائه شده است. بر اساس مقادیر α -برش بالایی و پایینی $\tilde{T}, \tilde{M}$ در سطوح مختلف، نمودار تقریبی ضریب همبستگی فازی بین مدیریت و تکنولوژی با استفاده از نرم

جدول ۱: مقادیر α -برش ضریب همبستگی فازی بین تکنولوژی و مدیریت شرکت‌های منتخب در نمونه تصادفی مثال (۱۰۲)

کرن	$\alpha = 0$	$\alpha = 1$	$\alpha = 0.2$	$\alpha = 0.3$	$\alpha = 0.4$	$\alpha = 0.5$	$\alpha = 0.6$	$\alpha = 0.7$	$\alpha = 0.8$	$\alpha = 0.9$	$\alpha = 1$
L	-۰/۸۲۸۱	-۰/۷۷۳۸	-۰/۶۹۳۳	-۰/۵۹۳۷	-۰/۴۸۶۱	-۰/۳۶۹۶	-۰/۲۴۲۶	-۰/۱۰۹۱	۰/۰۲۶۷	۰/۱۶۴۹	۰/۲۹۵۰
U	۰/۹۸۶۲	۰/۹۷۰۴	۰/۹۴۷۵	۰/۹۱۶۳	۰/۸۷۴۳	۰/۸۱۶۷	۰/۷۴۰۹	۰/۶۵۳۷	۰/۵۴۲۵	۰/۴۲۳۳	۰/۲۹۵۰

افزار متلب در شکل (۳) رسم شده است.



شکل ۳: نمودار تقریبی تابع عضویت ضریب همبستگی فازی بین تکنولوژی و مدیریت در مثال (۱۰۲)

برای تحلیل همبستگی بین این دو متغیر بر اساس شاخص ضریب همبستگی معرفی شده، ابتدا به نکات زیر که می‌توان از نتایج حاصل استخراج کرد توجه می‌کنیم (الف) تکیه‌گاه $\tilde{T}, \tilde{M}$ بسیار گسترده است و دامنه آن از 0.8281 تا 0.9862 می‌باشد. بر اساس این بازه به دست آمده، می‌توان چنین نتیجه‌گیری کرد که اگر چه ضریب همبستگی در این شرایط کمیته نادقیق است، اما نمی‌تواند کمتر از 0.8281 و بیشتر از 0.9862 باشد.

(ب) با فرض این که تمام داده‌های گزارش شده در این مثال اعدادی دقیق باشند ($\alpha = 1$) مقدار ضریب همبستگی بین دو متغیر تحت بررسی برابر با 0.2950 خواهد بود. از طرفی به سادگی می‌توان دید که میزان عضویت $r = 0.2950$ در مجموعه فازی به دست آمده برابر با ۱ است، علاوه بر آن می‌توانیم جهت استنباط دقیق تر میزان عضویت $r = 0$ که مبنی بر ناهمبسته بودن دو متغیر است را نیز بررسی کنیم. بر این اساس میزان عضویت $r = 0$ در مجموعه فازی به دست آمده برابر 0.78 است که مقدار قابل توجه است.

(پ) با توجه به نمودار به دست آمده برای تابع عضویت $\tilde{T}, \tilde{M}$ و با اندکی محاسبات، مساحت قرار گرفته در سمت چپ صفر برابر با $3/34$ درصد و در سمت راست آن $65/7$ درصد بوده است.

بر اساس سه مورد ذکر شده می‌توان در مجموع چنین نتیجه‌گیری کرد که ارتباط بین مدیریت و تکنولوژی در ۱۵ شرکت مورد بررسی نسبتاً ضعیف بوده است.

۳ رویکردی جدید برای ضریب همبستگی بین اعداد فازی تحت T -نرم دراستیک

همان‌طور که در بخش قبل گفته شد، نخستین رویکرد برای تعریف ضریب همبستگی بین اعداد فازی بر پایه استفاده از α -های اعداد فازی بود که منجر به حل یک برنامه‌ریزی ریاضی می‌شود. در این بخش با رویکردی متفاوت که بر پایه استفاده از T نرم‌ها بنا نهاده شده است، به بررسی تعریفی جدید برای ضریب همبستگی بین اعداد فازی می‌پردازیم.

۱.۳ عملیات جبری روی اعداد فازی بر اساس T -نرم دراستیک

در این بخش عملیات جبری چهارگانه شامل جمع، تفریق، ضرب، تقسیم بین اعداد فازی LR را بر اساس T -نرم دراستیک مورد مطالعه قرار می‌دهیم.

تعریف ۱.۳. فرض کنید $I = [0, 1]$. یک تابع دو متغیره به صورت $T : I \times I \rightarrow I$ را یک T -نرم^۱ گوییم اگر در خواص زیر صدق کند

$$T(x, 1) = x \text{ (الف)}$$

$$(ب) \text{ یکنوایی: } x_1 \leq x_2, y_1 \leq y_2 \implies T(x_1, y_1) \leq T(x_2, y_2)$$

$$(پ) \text{ جابه‌جایی: } T(x, y) = T(y, x)$$

$$(ت) \text{ شرکت‌پذیری: } T[x, T(y, z)] = T[T(x, y), z]$$

T -نرم‌ها و T -همنرم‌ها (یا S -نرم‌ها) را نرم‌های مثلثی و همنرم‌های مثلثی نیز گویند. این دو رده از اندازه‌ها به خاطر نقشی که در نامساوی‌های مثلثی دارند به این نام خوانده می‌شوند، به دلیل ویژگی‌های بسیار جالب اخیراً بیشتر مورد توجه واقع شده‌اند. در بین عملگرهای T -نرم عملگر دراستیک خواص جالبی دارد و مورد توجه قرار گرفته است و به‌صورت زیر تعریف می‌شود

$$T_W(x, y) = \begin{cases} x & y = 1, \\ y & x = 1, \\ 0 & x, y \neq 1. \end{cases}$$

فرض کنید $T = T_W$ ضعیف‌ترین T -نرم باشد و $\tilde{A} = (a, \alpha_A, \beta_A)_{LR}$ ، $\tilde{B} = (b, \alpha_B, \beta_B)_{LR}$ دو عدد فازی $L - R$ باشند. در این صورت جمع بین این دو عدد فازی تحت T -نرم دراستیک عبارتست از

$$\tilde{A} \oplus \tilde{B} = (a, \alpha_A, \beta_A)_{LR} \oplus (b, \alpha_B, \beta_B)_{LR} = (a + b, \max(\alpha_A, \alpha_B), \max(\beta_A, \beta_B))_{LR}. \quad (11)$$

در صورتی که $L = R$ باشد. تفاضل این دو عدد فازی تحت T -نرم دراستیک یک عدد فازی LR به‌صورت زیر است

$$\begin{aligned} \tilde{A} \ominus \tilde{B} &= (a, \alpha_A, \beta_A)_{LR} \ominus (b, \alpha_B, \beta_B)_{LR} \\ &= (a, \alpha_A, \beta_A)_{LR} \oplus (-b, \beta_B, \alpha_B)_{LR} = (a - b, \max(\alpha_A, \beta_B), \max(\beta_A, \alpha_B))_{LR}. \end{aligned} \quad (12)$$

برای بررسی ضرب دو عدد فازی LR تحت T -نرم دراستیک بایستی حالت‌های مختلفی را برای مراکز دو عدد فازی تحت بررسی مورد مطالعه قرار دهیم که در ادامه به آن پرداخته شده‌است.

حالت اول: فرض کنید $a, b > \circ$ در این صورت برای مقادیر $ab \geq z$ داریم

$$\begin{aligned} (\tilde{A} \otimes \tilde{B})(z) &= \sup_{x, y=z} T_W(\tilde{A}(x), \tilde{B}(y)) \\ &= \max(\tilde{A}(\frac{z}{b}), \tilde{B}(\frac{z}{a})) = \max(L(\frac{a - \frac{z}{b}}{\alpha_A}), L(\frac{b - \frac{z}{a}}{\alpha_B})) \\ &= \max(L(\frac{ab - z}{\alpha_A b}), L(\frac{ab - z}{\alpha_B a})), \end{aligned} \quad (13)$$

بنابراین می توان نوشت

$$(\tilde{A} \otimes \tilde{B})(z) = \begin{cases} L[\frac{(ab-z)}{\max(\alpha_A b, \alpha_B a)}] & z \geq ab - \max(\alpha_A b, \alpha_B a), \\ \circ & \text{در غیر این صورت} \end{cases} \quad (14)$$

به طور مشابه

$$(\tilde{A} \otimes \tilde{B})(z) = \begin{cases} R[\frac{(z-ab)}{\max(\beta_A b, \beta_B a)}] & z \leq ab - \max(\beta_A b, \beta_B a), \\ \circ & \text{در غیر این صورت} \end{cases} \quad (15)$$

با جمع بندی دو عبارت به دست آمده در روابط (14) و (15) می توان ضرب دو عدد فازی LR را در حالت $a, b > \circ$ به صورت یک عدد فازی LR با تعریف زیر ارائه نمود

$$\tilde{A} \otimes \tilde{B} = (ab, \max(\alpha_A b, \alpha_B a), \max(\beta_A b, \beta_B a))_{LR}. \quad (16)$$

به روش مشابه با آنچه که در حالت اول مورد بررسی قرار گرفت، حالت های ممکن دیگر برای a, b به عنوان مراکز دو عدد فازی تحت بررسی در ادامه آورده شده است. حالت دوم: اگر $a, b < \circ$ آن گاه ضرب دو عدد فازی $\tilde{A}$ و $\tilde{B}$ یک عدد فازی LR با تعریف زیر می باشد

$$\tilde{A} \otimes \tilde{B} = (ab, \max(\beta_A b, \beta_B a), \max(\alpha_A b, \alpha_B a))_{RL}.$$

حالت سوم: فرض کنید $a = \circ, b > \circ$ در این صورت حاصل ضرب دو عدد فازی $\tilde{A}$ و $\tilde{B}$ عدد فازی LR با تعریف زیر است

$$\tilde{A} \otimes \tilde{B} = (\circ, \alpha_A b, \beta_A b)_{LR}.$$

حالت چهارم: برای حالتی که $a = \circ, b < \circ$ باشد، حاصل ضرب یک عدد فازی LR به صورت زیر است

$$\tilde{A} \otimes \tilde{B} = (\circ, -\beta_A b, -\alpha_A b)_{RL}.$$

حالت پنجم: در صورتی که $a = \circ, b = \circ$ باشند آنگاه

$$\tilde{A} \otimes \tilde{B} = (\circ, \circ, \circ)_{LR} = I_{\{\circ\}}(x).$$

که در آن

$$I_{\{\circ\}}(x) = \begin{cases} 1 & x = \circ \\ \circ & \text{در غیر این صورت} \end{cases}.$$

حالت ششم: در این حالت فرض می‌کنیم $a < 0, b > 0$ باشند و علاوه بر آن $L = R$ باشد. در این صورت

$$\tilde{A} \otimes \tilde{B} = (ab, \max(\alpha_A b, -\beta_B a), \max(\beta_A b, -\alpha_B a))_{RR}.$$

در صورتی که $\tilde{A}$ و $\tilde{B}$ اعداد فازی متقارن باشند آنگاه حاصل ضرب این دو عدد فازی متقارن یک عدد فازی LR با تعریف زیر می‌باشد

$$(\tilde{A} \otimes \tilde{B})(z) = (ab, \max(\alpha_A |b|, \alpha_B |a|), \max(\alpha_A |b|, \alpha_B |a|))_{LL}.$$

برای به دست آوردن تابع عضویت ضریب همبستگی فازی $\tilde{r}_{\tilde{X}, \tilde{Y}}$ نیازمند محاسبه تابع عضویت آماره‌های میانگین و انحراف می‌باشیم. در ادامه بدون کاستن از کلیت مسئله فرض می‌کنیم که $\tilde{A}, \tilde{B}$ دو عدد فازی مثلثی متقارن باشند. در این بخش فرض می‌کنیم که $\tilde{A} = (a, \alpha)$ اگر $\tilde{A} + \dots + \tilde{A} = n \odot \tilde{A}$ باشد این صورت اگر $\tilde{A} = (a, \alpha)$ باشد آنگاه $n \odot \tilde{A} = \tilde{A} + \dots + \tilde{A} = (na, \alpha)$. بنابراین برای هر مقدار غیر صفر مانند t داریم $t \odot \tilde{A} = (ta, \alpha)$. اکنون فرض کنید که n مشاهده فازی از متغیرهای Y, X در اختیار داشته باشیم و هدف محاسبه ضریب همبستگی این زوج مشاهدات نادقیق باشد. در چنین شرایطی رابطه (۲) را بر حسب نماد گذاری‌های جدید انجام شده، به صورت زیر بازنویسی می‌کنیم (۲)

$$\tilde{r}_{\tilde{X}, \tilde{Y}} = \frac{\sum_{j=1}^n (\tilde{X}_j - \frac{1}{n} \odot \sum_{j=1}^n \tilde{X}_j)(\tilde{Y}_j - \frac{1}{n} \odot \sum_{j=1}^n \tilde{Y}_j)}{\sqrt{\sum_{j=1}^n (\tilde{X}_j - \frac{1}{n} \odot \sum_{j=1}^n \tilde{X}_j)^2 \sum_{j=1}^n (\tilde{Y}_j - \frac{1}{n} \odot \sum_{j=1}^n \tilde{Y}_j)^2}}. \quad (17)$$

فرض کنید $\tilde{Y}_j = (y_j, \delta_j)_L, \tilde{X}_j = (x_j, r_j)_L$ نمونه‌های تصادفی از جوامع تحت بررسی می‌باشند که اعداد فازی متقارن در نظر گرفته شده‌اند. $(j = 1, \dots, n)$ با توجه به عملگرهای جبری تعریف شده بر اساس $-T$ نرم دراستیک و نمادگذاری‌های انجام شده می‌توان نوشت

$$\tilde{\tilde{X}} = \frac{1}{n} \odot \sum_{j=1}^n \tilde{X}_j = (\frac{1}{n} \sum_{j=1}^n x_j, \max_{1 \leq j \leq n} r_j)_L, \quad \tilde{\tilde{Y}} = \frac{1}{n} \odot \sum_{j=1}^n \tilde{Y}_j = (\frac{1}{n} \sum_{j=1}^n y_j, \max_{1 \leq j \leq n} \delta_j)_L$$

در این صورت با استفاده از رابطه (۱۶) خواهیم داشت

$$(\tilde{X}_j - \tilde{\tilde{X}})(\tilde{Y}_j - \tilde{\tilde{Y}}) = ((x_j - \frac{1}{n} \sum_{j=1}^n x_j)(y_j - \frac{1}{n} \sum_{j=1}^n y_j), \max(|x_j - \frac{1}{n} \sum_{j=1}^n x_j|, \max_{1 \leq j \leq n} \delta_j, |y_j - \frac{1}{n} \sum_{j=1}^n y_j|, \max_{1 \leq j \leq n} r_j))_L.$$

با استفاده از رابطه (۱۱) داریم

$$\begin{aligned} \sum_{j=1}^n (\tilde{X}_j - \frac{1}{n} \odot \sum_{j=1}^n \tilde{X}_j)(\tilde{Y}_j - \frac{1}{n} \odot \sum_{j=1}^n \tilde{Y}_j) = \\ (\sum_{j=1}^n (x_j - \frac{1}{n} \sum_{k=1}^n x_k)(y_j - \frac{1}{n} \sum_{k=1}^n y_k), \max_{1 \leq j \leq n} (|x_j - \frac{1}{n} \sum_{k=1}^n x_k|, \max_{1 \leq k \leq n} \delta_k, |y_j - \frac{1}{n} \sum_{k=1}^n y_k|, \max_{1 \leq k \leq n} r_k))_L. \end{aligned} \quad (18)$$

از طرفی به طور مشابه می‌توان نوشت

$$\sum_{j=1}^n (\tilde{X}_j - \frac{1}{n} \odot \sum_{j=1}^n \tilde{X}_j)^2 = (\sum_{j=1}^n (x_j - \frac{1}{n} \sum_{k=1}^n x_k)^2, \max_{1 \leq j \leq n} |x_j - \frac{1}{n} \sum_{k=1}^n x_k|, \max_{1 \leq k \leq n} r_k)_L$$

بنابراین

$$\begin{aligned}
 & \sqrt{\sum_{j=1}^n (\tilde{X}_j - \frac{1}{n} \odot \sum_{k=1}^n \tilde{X}_k)^2 \sum_{j=1}^n (\tilde{Y}_j - \frac{1}{n} \odot \sum_{k=1}^n \tilde{Y}_k)^2} \\
 &= \sqrt{\sum_{j=1}^n (x_j - \frac{1}{n} \sum_{k=1}^n x_k)^2 \sum_{j=1}^n (y_j - \frac{1}{n} \sum_{k=1}^n y_k)^2} \\
 & \quad , \sqrt{\sum_{j=1}^n (x_j - \frac{1}{n} \sum_{k=1}^n x_k)^2 \sum_{j=1}^n (y_j - \frac{1}{n} \sum_{k=1}^n y_k)^2} \\
 & \times \max\{\sum_{j=1}^n (x_j - \frac{1}{n} \sum_{k=1}^n x_k)^2 \max_{1 \leq j \leq n} |y_j - \frac{1}{n} \sum_{k=1}^n y_k| \max_{1 \leq k \leq n} \delta_k, \\
 & \sum_{j=1}^n (y_j - \frac{1}{n} \sum_{k=1}^n y_k)^2 \max_{1 \leq j \leq n} |x_j - \frac{1}{n} \sum_{k=1}^n x_k| \max_{1 \leq k \leq n} r_k\} \}_L.
 \end{aligned} \tag{۱۹}$$

در نهایت می‌توان تابع عضویت ضریب همبستگی فازی $\tilde{r}$ را با استفاده از روابط به‌دست آمده برای تقسیم اعداد فازی تحت T - نرم دراستیک محاسبه نمود.

مثال ۲.۳. یک مجموعه داده شامل یک زوج مشاهده از اعداد فازی مثلثی متقارن که در جدول (۹۹) آورده شده‌اند را در نظر بگیرید (۱)

جدول ۲: مجموعه داده شامل ۸ جفت از اعداد فازی مثلثی متقارن در مثال (۲.۳)

i	$\tilde{X}_i$	$\tilde{Y}_i$
۱	(۲/۰, ۰/۵)	(۴/۰, ۰/۵)
۲	(۳/۵, ۰/۵)	(۵/۵, ۰/۵)
۳	(۵/۵, ۱/۰)	(۷/۵, ۱/۰)
۴	(۷/۰, ۰/۵)	(۶/۵, ۰/۵)
۵	(۸/۵, ۰/۵)	(۸/۵, ۵/۰)
۶	(۱۰/۵, ۱/۰)	(۸/۰, ۱/۰)
۷	(۱۱/۰, ۰/۵)	(۱۰/۵, ۰/۵)
۸	(۱۲/۵, ۰/۵)	(۹/۵, ۰/۵)

برای محاسبه ضریب همبستگی با استفاده از رابطه (۹۹) می‌توان نوشت

$$\sum_{j=1}^n (\tilde{X}_j - \tilde{\bar{X}})(\tilde{Y}_j - \tilde{\bar{Y}}) = (۵/۷۵۰۰, ۵/۵۶۲۵)$$

بنابراین

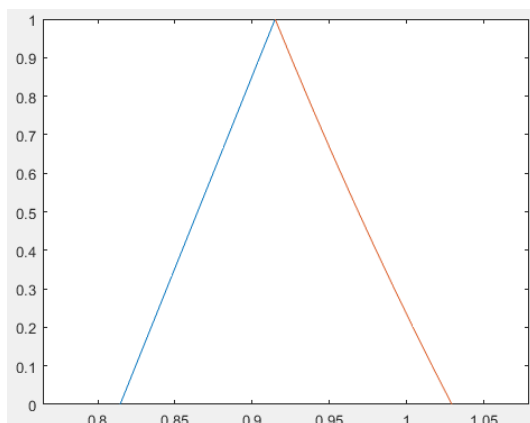
$$\sqrt{\sum_{j=1}^n (\tilde{X}_j - \tilde{\bar{X}})^2 \sum_{j=1}^n (\tilde{Y}_j - \tilde{\bar{Y}})^2} = (۵۵/۴۸۱۰, ۶/۱۶۴۶)$$

در این صورت ضریب همبستگی اعداد فازی برابر است

$$\tilde{r}_{\tilde{X}, \tilde{Y}} = \frac{(\tilde{0}/\tilde{7500}, \tilde{0}/\tilde{5625})}{(\tilde{55}/\tilde{4810}, \tilde{6}/\tilde{1646})} = \begin{cases} 1 - \frac{\tilde{0}/\tilde{9147} - z}{\tilde{0}/\tilde{180} \times \max(\tilde{0}/\tilde{5625}, \tilde{6}/\tilde{1646} \times z)} & \tilde{0}/\tilde{8145} \leq z \leq \tilde{0}/\tilde{9147}, \\ 1 - \frac{z - \tilde{0}/\tilde{9147}}{\tilde{0}/\tilde{180} \times \max(\tilde{0}/\tilde{5625}, \tilde{6}/\tilde{1646} \times z)} & \tilde{0}/\tilde{9147} \leq z \leq \tilde{1}/\tilde{0291}, \\ \circ & \text{در غیر این صورت.} \end{cases} \quad (20)$$

$$= \begin{cases} 1 - \frac{\tilde{0}/\tilde{9147} - z}{\tilde{0}/\tilde{1000}} & \tilde{0}/\tilde{8145} \leq z \leq \tilde{0}/\tilde{9023}, \\ 1 - \frac{\tilde{0}/\tilde{9147} - z}{\tilde{0}/\tilde{111} \times z} & \tilde{0}/\tilde{9023} \leq z \leq \tilde{0}/\tilde{9147}, \\ 1 - \frac{z - \tilde{0}/\tilde{9147}}{\tilde{0}/\tilde{111} \times z} & \tilde{0}/\tilde{9147} \leq z \leq \tilde{1}/\tilde{0291}, \\ \circ & \text{در غیر این صورت.} \end{cases} \quad (21)$$

نمودار مربوط به تابع عضویت ضرب همبستگی فازی مثلثی متقارن در شکل زیر نمایش داده شده است.



شکل ۴: نمودار تقریبی تابع عضویت ضرب همبستگی فازی مثلثی متقارن در مثال (۲.۳)

مراجع

- [1] Chachi, J., Taheri, S.M., (2016), Multiple Fuzzy Regression Model for Fuzzy Input Output Data, Iranian Journal of Fuzzy Systems, Vol 13, pp 63-78.
- [2] Hong D. H., (2006), "Fuzzy measures for a correlation coefficient of fuzzy numbers under T_W (the weakest t -norm)-based fuzzy arithmetic operations", Information Sciences, Vol 176, pp 150-160.
- [3] Liu, S. T., Kao, C., (2002), "Fuzzy measures for correlation coefficient of fuzzy numbers", Fuzzy Sets and Systems, Vol 128, pp 267-275.

مشارکت‌های استاد ارقامی در زمینه آمار و احتمال فازی

سید محمود طاهری^۱

دانشگاه تهران، دانشکده فنی

چکیده

نخستین مطالعات در حوزه آمار و احتمال فازی در ایران به حدود ربع قرن قبل (ابتدای دهه هفتاد شمسی) باز می‌گردد. استاد مرحوم دکتر ناصر رضا ارقامی نقش ویژه‌ای در شروع و تداوم این مطالعات داشتند. در این مختصر، تاثیر ایشان را در تحقیقات علمی مرتبط با آمار و احتمال فازی شرح می‌دهیم، به مقالات ایشان در این حوزه اشاره می‌کنیم و همچنین به نقش استاد ارقامی در زمینه‌های آموزشی و ترویجی آمار و احتمال فازی می‌پردازیم. سرانجام به فعالیت‌های علمی و اجرایی ایشان در انجمن سیستم‌های فازی ایران اشاره می‌کنیم.

کلمات کلیدی: آمار فازی، احتمال فازی، رگرسیون فازی

فرصت ز رنگ ماست پرافشان نیستی غافل ز ما مباش، که ناگاه رفته‌ایم (بیدل دهلوی)

۱ مقدمه

نظریه و کاربرد سیستم‌های فازی، به وسیله دانشمند ایرانی‌تبار مرحوم پروفسور لطفی عسگرزاده، نامدار به زاده، (۱۲۹۹-۱۳۹۶ ه.ش، ۱۹۲۱-۲۰۱۷ م) در سال ۱۳۴۳ ه.ش. (۱۹۶۵ م) معرفی شد، (۱۸).^۲ این نظریه سال‌ها در ایران تقریباً ناشناخته بود. توجه جدی به این نظریه در ایران، به دهه شصت (ه.ش.) باز می‌گردد. در آن زمان دو جریان علمی در دانشگاه تهران (دانشکده فنی) و دانشگاه شهید باهنر کرمان پا گرفت، که اولی بیشتر در حوزه‌های فنی-مهندسی و دومی بیشتر در حوزه‌های ریاضیات و منطق فازی بود. البته در برخی دانشگاه‌های دیگر نیز مانند دانشگاه صنعتی اصفهان، دانشگاه شیراز و دانشگاه صنعتی شریف فعالیت‌های علمی توسط برخی

^۱ sm-taheri@ut.ac.ir

^۲ برای "سال‌شمار زندگی پروفسور زاده" و "زندگی‌نامه کوتاه پروفسور زاده" به سایت انجمن سیستم‌های فازی ایران به نشانی www.fuzzy.ir مراجعه کنید. نیز خبرنامه پاییز و زمستان ۱۳۹۴ انجمن سیستم‌های فازی ایران (به همان نشانی) را ببینید.

اساتید در حوزه سیستم‌های فازی آغاز شده بود. در این زمینه تعدادی مقاله مروری/ترویجی/ترجمه‌ای نیز منتشر شده بود و چند پایان‌نامه کارشناسی‌ارشد به رشته تحریر درآمده بود.

مطالعات در زمینه آمار و احتمال فازی در ایران، از اوایل دهه هفتاد شروع شد. در این یادداشت کوتاه، به نقش استاد مرحوم دکتر ارقامی در این زمینه می‌پردازیم. آنچه در پی می‌آید، مبتنی بر مستندات در دسترس و چاپ‌شده است. آشکارست که نقش غیرمستقیم و پوشیده‌استاد ارقامی در بسیاری از زمینه‌ها از جمله آمار و احتمال فازی بیش از اینهاست.

۲ مقالات و سخنرانی‌ها در همایش‌ها

نخستین سخنرانی دکتر ارقامی در زمینه آمار و احتمال فازی، در همایش زیر صورت گرفت:

نخستین سمینار مجموعه‌های مشکک و کاربرد آن (دانشگاه شهید باهنر کرمان، ۱۳۷۲)

عنوان سخنرانی، ”مروری بر رگرسیون فازی“ بود. متن مقاله ایشان که به‌صورت دست‌نویس و هجده صفحه بود در گزارش سمینار چاپ شد، (۱). سمینار فوق، شاید اولین همایشی بود که در آن جمع چشم‌گیری از اساتید و پژوهشگران هر دو حوزه نظری و کاربردی سیستم‌های فازی شرکت داشتند. از جمله بجاست به اساتید مرحوم دکتر کارولوکس و دکتر محمدعلی پورعبدا... اشاره کنم. گفتنی است که در حاشیه سمینار نشستی در باره معادله‌یابی برای واژه فازی برگزار شد. در نشست، دکتر پورعبدا... واژه شولا، دکتر ارقامی واژه تار، و دکتر ماشین‌چی واژه مشکک را به‌عنوان معادله‌هایی برای واژه فازی پیشنهاد دادند. این سخنرانی، به‌صورت بروزشده، در دومین کارگاه آمار و احتمال فازی (دانشگاه صنعتی اصفهان، ۲۳ و ۲۴ اسفند ۱۳۸۶) مجدداً توسط ایشان ارائه گردید و سپس نسخه تکمیل‌شده آن تهیه شد و در شماره ۲۳ اندیشه آماری (نشریه انجمن آمار ایران) به چاپ رسید، (۱۲).

همچنین، ایشان در چهار مقاله دیگر کنفرانسی مشارکت داشتند. اولین آنها، مقاله (۹) بود که در نخستین کنفرانس آمار ایران ارائه شد. در این مقاله، موضوع احتمال پیشامدهای فازی طبق رویکرد اول بیگر، رویکرد کلمنت و رویکرد دوم بیگر بررسی شده است. دیگری، مقاله (۳) بود که در آن یک قضیه حد مرکزی برای متغیرهای تصادفی فازی بررسی شده است. سه دیگر، مقاله (۶) است که در این مقاله رگرسیون فازی بازه‌ای-مقدار مبتنی بر ملاک کمترین مربعات خطا بررسی شده است. و سرانجام، به مقاله (۱۴) اشاره می‌کنیم که در کنفرانس تخصصی IEEE در سیستم‌های فازی ارائه شد و در آن یک مدل رگرسیون فازی بازه‌ای-مقدار معرفی شده است.

۳ راهنمایی پایان‌نامه‌های کارشناسی‌ارشد و رساله‌های دکتر

مرحوم دکتر ارقامی راهنمایی دو پایان‌نامه و یک رساله را در حوزه آمار و احتمال فازی به‌عهده داشتند. اولین پایان‌نامه کارشناسی‌ارشد که تحت راهنمایی ایشان انجام شد (به احتمال زیاد، نخستین پایان‌نامه کارشناسی‌ارشد آمار در موضوع آمار و احتمال فازی در ایران) در سال ۱۳۷۰ بود، (۸) (استاد مشاور: مرحوم دکتر محمدعلی پورعبدا...). در این پایان‌نامه، مطالب زیر معرفی/بررسی/تحلیل شده است: اندازه‌های احتمال پیشامدهای فازی، اندازه احتمال فازی، متغیر تصادفی فازی، محاسبه احتمالات براساس مشاهدات فازی، آزمون فرضیه‌های فازی (رویکردهای کلاسیک و بیز) و نظریه امکان. دیگر، پایان‌نامه کارشناسی‌ارشد آقای سعید اخلاقی بود، (۲)، که در آن مدل رگرسیون فازی، براساس رویکرد مبتنی بر کمترین مربعات خطا، بررسی و تحلیل شده است. به‌علاوه، ایشان استاد راهنمای رساله دکترای آقای دکتر محمدرضا ربیعی بودند، (۷). در این رساله، رگرسیون کمترین مربعات برای داده‌های ورودی-خروجی فازی بازه‌ای-مقدار، رگرسیون کمترین مربعات در محیط تماماً فازی بازه‌ای-مقدار، و همچنین رگرسیون امکانی در محیط فازی بازه‌ای-مقدار بررسی و تحلیل شده است.

استاد ارقامی استاد مشاور دو رساله دکتر در زمینه آمار و احتمال فازی بودند. اولین آنها رساله آقای دکتر محسن عارفی بود، (۱۱)، که در آن مباحثی مانند آزمون فرضیه‌های فازی براساس آماره آزمون فازی (دو حالت داده‌های دقیق و داده‌های نادقیق)، آزمون فرضیه‌های فازی براساس فاصله اطمینان فازی، برآورد بیز امکانی (در دو حالت داده‌های دقیق و داده‌های فازی) و برآورد تابع چگالی فازی براساس داده‌های فازی، بررسی شد. دومین رساله دکتر در حوزه آمار و احتمال فازی که مرحوم استاد ارقامی استاد مشاور بودند رساله آقای دکتر جلال چاچی بود، (۵). در این رساله، تشکیل فواصل اطمینان فازی برای میانگین متغیرهای تصادفی فازی، آزمون فرضیه

در محیط فازی (رویکرد مبتنی بر فواصل اطمینان)، پرتوان‌ترین آزمون در محیط تماماً فازی، و همچنین چند موضوع در رگرسیون فازی از جمله رگرسیون استوار فازی با پهناهای متغیر، رگرسیون فازی کمترین مربعات خطا برای مدل با ورودی-خروجی فازی، و رگرسیون فازی کمترین قدرمطلق خطا براساس متر هاسدورف تعمیم‌یافته، بررسی شده است.

۴ مقاله‌های علمی

مقاله‌های پژوهشی استاد مرحوم دکتر ارقامی در حوزه آمار و احتمال فازی، شامل چهار مقاله، بر محور رگرسیون فازی متمرکز بوده است. نخستین آنها، مقاله (۱۵) بود. در این مقاله، مدل‌سازی رگرسیونی در محیط تماماً فازی بازه‌ای-مقدار، طبق روش کمترین مربعات تعمیم‌یافته، انجام گرفته است. منظور از محیط تماماً فازی بازه‌ای-مقدار این است که هر سه مولفه‌های مدل یعنی داده‌های متغیر(های) مستقل، داده‌های متغیر وابسته و ضرائب مدل، اعداد فازی بازه‌ای-مقدار هستند. سپس مقاله (۱۳) بود که در این مقاله، یک مدل رگرسیون ترکیبی (مبتنی بر یک شیوه بهینه‌سازی و رگرسیون مارس (MARS)) برای مدل‌سازی و تحلیل داده‌های فازی معرفی شد و کاربرد آن در مهندسی آب (هیدرولوژی) براساس مجموعه داده‌های واقعی از رواناب‌های ثبت‌شده در ناحیه‌ای از خراسان شمالی بررسی شد. روشی که در این مقاله معرفی شد، همزمان دو مشکل تغییرات زیاد در داده‌ها و مشکل پهناهای (ابهام‌های) متغیر در داده‌های فازی را حل‌وفصل می‌کند. سه دیگر، مقاله (۱۵) است که در این مقاله، مدل‌سازی رگرسیونی برای داده‌های فازی بازه‌ای-مقدار، مبتنی بر رویکرد امکانی انجام گرفته است. در این رویکرد، مسئله برآورد ضرائب مدل، به یک مسئله برنامه‌ریزی خطی تبدیل می‌شود و با حل آن، ضرائب مدل که اعداد فازی بازه‌ای-مقدار هستند برآورد می‌شوند. و سرانجام، به یک مقاله در دست داوری اشاره می‌شود که مربوط است به مدل‌سازی رگرسیونی براساس داده‌های ورودی-خروجی فازی بازه‌ای-مقدار، مبتنی بر روش کمترین مربعات خطای تعمیم‌یافته طبق چندین معیار فاصله برای اعداد فازی بازه‌ای-مقدار، و کاربردهای آن در دو حوزه آب‌شناسی و خاک‌شناسی، (۱۷).

بجاست به دو مقاله دیگر از مرحوم استاد ارقامی اشاره شود که در واقع اولین و آخرین مقاله (غیرکنفرانسی) بوده است که ایشان مشارکت داشته‌اند. هر دو این مقالات، ترویجی بوده‌اند. اولین مقاله در باره تشریح و تبیین مفهوم و کاربرد متغیرهای تصادفی فازی بوده است، (۱۰). آخرین مقاله چاپ‌شده در زمینه آمار و احتمال فازی که استاد ارقامی در آن مشارکت داشتند (و در آخرین ماه‌های حیات ایشان چاپ شد)، مقاله‌ای بود که در آن به طرز مفهومی و با استفاده از مثال‌های عددی/کاربردی، سه رویکرد کلاسیک (رویکرد نیمن-پیرسونی)، رویکرد بیز و رویکرد کم-بیشینه (مینیمکس) به مسئله آزمون فرضیه‌های فازی تبیین و تشریح شده است، (۴).

۵ فعالیت‌های علمی-اجرایی

دکتر ارقامی، از اعضای هیات مؤسس مجله پژوهشی انجمن سیستم‌های فازی ایران (IJFS) Systems Fuzzy of Journal Iraninan بودند. گفتنی است، در ۲۹ و ۳۰ خرداد ۱۳۸۱، "سومین همایش مجموعه‌های فازی و کاربردهای آن"، در دانشگاه سیستان و بلوچستان (شهر زاهدان) برگزار شد. در حاشیه این همایش، طی جلسه‌ای طولانی، راه‌اندازی یک مجله علمی-پژوهشی در زمینه سیستم‌های فازی مورد بحث و بررسی قرار گرفت و نسخه اولیه اساس‌نامه مجله تهیه شد. این اساس‌نامه، با تغییراتی اندک (مصوب ۱۳۸۱/۹/۲۱ هیأت مؤسس)، مبنای چاپ و انتشار مجله IJFS شد. همچنین، مرحوم دکتر ارقامی از ابتدای انتشار مجله (سال ۱۳۸۴ ه.ش، ۲۰۰۴ م.) عضو هیات تحریریه این مجله بودند.

دکتر ارقامی، از اعضای هیات مؤسس انجمن سیستم‌های فازی ایران بودند. شایان ذکر است که در پی برگزاری چندین کارگاه، سمینار و کنفرانس در اواخر دهه شصت و دهه هفتاد (ه.ش.) و سپس تأسیس شاخه سیستم‌های فازی در انجمن آمار ایران در سال ۱۳۸۱، موضوع تأسیس یک انجمن علمی مستقل در زمینه سیستم‌های فازی، در حاشیه "چهارمین کنفرانس مجموعه‌های فازی و کاربردهای آن" (دانشگاه مازندران - شهر بابل) در شهریور ۱۳۸۲، مطرح و پیش‌نویس اساس‌نامه توسط هیأت مؤسس تدوین شد. سرانجام انجمن سیستم‌های فازی ایران با تصویب اساس‌نامه آن در ۱۳۸۳/۷/۲۵ در کمیسیون انجمن‌های علمی وزارت علوم رسماً تأسیس شد. مرحوم دکتر ارقامی در نخستین مجمع عمومی انجمن مورخ ۱۳۸۴/۳/۵ به‌عنوان عضو هیات مدیره انجمن (به مدت دو سال) انتخاب شدند.

در پایان بجاست اشاره شود که استاد ارقامی در برنامه‌ریزی و پشتیبانی از برگزاری مجموعه سخنرانی‌هایی در زمینه ریاضیات فازی و احتمال و آمار فازی در سال

سپاسگزاری

از برگزارکنندگان هشتمین سمینار آمار و احتمال فازی (دانشکده علوم ریاضی دانشگاه فردوسی مشهد) که مراسم یادبود مرحوم استاد ارقامی را در برنامه سمینار منظور نمودند سپاسگزاری می‌کنم. همان‌گونه که در ابتدا اشاره شد، مطالب این یادداشت، مبتنی بر اطلاعات و مستندات نویسنده بوده است. بسا مواردی باشد که نویسنده از آنها اطلاع نداشته است. از علاقه‌مندان درخواست می‌شود چنانچه کاستی/اشتباه در مطالب و مستندات دیدند به نگارنده یادآور شوند.

مراجع

[۱] ارقامی، ن.ر. (۱۳۷۲). مروری بر رگرسیون فازی، گزارش نخستین سمینار مجموعه‌های مشکک و کاربرد آن، دانشگاه شهید باهنر کرمان، مقاله شماره ۴۵، صص ۱۸-۱.

[۲] اخلاقی، س. (۱۳۸۳). رگرسیون حداقل مربعات فازی، پایان‌نامه کارشناسی ارشد آمار ریاضی، دانشگاه فردوسی مشهد (دانشکده علوم ریاضی).

[۳] اکبری، م.ق. و ارقامی، ن.ر. (۱۳۸۴). قضیه حد مرکزی برای متغیرهای تصادفی فازی، گزارش پنجمین سمینار احتمال و فرایندهای تصادفی، دانشگاه بیرجند، مقاله شماره ۴۵، صص ۱۸-۱.

[۴] پرچمی، ع.، طاهری، س.م. و ارقامی، ن.ر. (۱۳۹۵). مقایسه‌ای از رویکردهای نیمن-پیرسون، بیز و کم-بیشینه در آزمون فرضیه‌های فازی، فرهنگ و اندیشه ریاضی، ۵۹، ۴۱-۶۸.

[۵] چاچی، ج. (۱۳۹۱). روش‌های آماری براساس اطلاعات نادقیق، رساله دکترای آمار ریاضی، دانشگاه صنعتی اصفهان (دانشکده علوم ریاضی).

[۶] ربیعی، م.ر.، ارقامی، ن.ر.، طاهری، س.م. و صادق‌پور، ب. (۱۳۹۱). رگرسیون کمترین توان‌های دوم برای داده‌های ورودی و خروجی فازی بازه‌ای-مقدار، گزارش یازدهمین کنفرانس سیستم‌های هوشمند ایران، دانشگاه خوارزمی، مقاله شماره ۴۵، صص ۸-۱.

[۷] ربیعی، م.ر. (۱۳۹۲). مدل‌سازی رگرسیونی در محیط فازی بازه‌ای-مقدار، رساله دکترای آمار ریاضی، دانشگاه فردوسی مشهد (دانشکده علوم ریاضی).

[۸] طاهری، س.م. (۱۳۷۰). کاربرد نظریه مجموعه‌های فازی در آمار و احتمالات، پایان‌نامه کارشناسی ارشد آمار ریاضی، دانشگاه فردوسی مشهد (دانشکده علوم، گروه آمار).

[۹] طاهری، س.م. و ارقامی، ن.ر. (۱۳۷۱). احتمال پیشامدهای فازی، مجموعه مقالات نخستین کنفرانس آمار ایران، دانشگاه صنعتی اصفهان، صص ۲۹۷-۳۳۰.

[۱۰] طاهری، س.م. و ارقامی، ن.ر. (۱۳۷۴). متغیرهای تصادفی فازی، گلچین ریاضی (نشریه بخش ریاضی دانشگاه شیراز)، ۳، ۶-۱۵.

[۱۱] عارفی، م. (۱۳۸۹). استنباط آماری در محیط فازی، رساله دکترای آمار ریاضی، دانشگاه صنعتی اصفهان (دانشکده علوم ریاضی).

[۱۲] میرزایی یگانه، ش. و ارقامی، ن.ر. (۱۳۸۶). رگرسیون فازی: مروری بر چند رویکرد، اندیشه آماری، ۲۳، ۳۵-۴۷.

- [14] Rabiee M.R., Arghami, N.R., Taheri, S.M. and Sadeghpour, B. (2013), Fuzzy regression model with interval-valued fuzzy input-output data, *Proceedings of the 2013 IEEE International Conference on Fuzzy Systems*, Heydarabad (India), pp. 201-206.
- [15] Rabiee M.R., Arghami, N.R., Taheri, S.M. and Sadeghpour, B. (2014), Least-squares approach to regression modeling in full interval-valued fuzzy environment, *Soft Computing* 18, 2043-2059.
- [16] Rabiee M.R., Taheri, S.M. and Arghami, N.R. (2015), A linearprogramming approach to interval-valued fuzzy regression analysis, *International Journal of Intelligent Technologies and Applied Statistics* 8, 171-203.
- [17] Rabiee M.R., Taheri, S.M. and Arghami, N.R. and Sadeghpour, B. (2016), Least-squares regression analysis for interval-valued fuzzy input-output data, *Under review*.
- [18] Zadeh, L.A. (1965), Fuzzy sets, *Information and Control*, 8, 338-353.

فازی سازی داده های باستان‌شناسی به منظور استفاده در آمار فازی

سید محمود طاهری^۱، فرشید ایروانی قدیم^۲، محمدعلی کبیریان^۳

^۱ دانشگاه تهران، دانشکده فنی

^۲ دانشگاه هنر اصفهان، دانشکده حفاظت و مرمت

^۳ دانشگاه هنر اصفهان، دانشکده حفاظت و مرمت

چکیده

این بررسی نگرشی بر استفاده از ریاضیات فازی و آمار فازی در باستان‌شناسی است. به همین منظور ابتدا لزوم استفاده از ریاضیات و آمار فازی در باستان‌شناسی مورد بحث قرار می‌گیرد. سپس به موضوع فازی‌سازی داده‌های باستان‌شناسی به منظور استفاده از روش‌های آمار فازی در باستان‌شناسی پرداخته می‌شود. همچنین تشکیل جدول فراوانی براساس داده‌های فازی و مزیت آن تشریح می‌شود. مطالب با چندین مثال کاربردی در حوزه‌ی باستان‌شناسی تبیین شده‌اند.

کلمات کلیدی: باستان‌شناسی، فازی‌سازی، آمار فازی

۱ مقدمه

”باستان‌شناسی علم مطالعه‌ی گذشته با استفاده از مستندات تجربی است” (۲). در باستان‌شناسی برای تجزیه و تحلیل داده‌ها و مستندات بدست آمده از کاوش، از روش‌های گوناگونی مانند آزمایش‌های شیمیایی، تحلیل‌های هنری، زیست‌شناسی و ریاضیات و آمار استفاده می‌شود. استفاده از روش‌های آماری و ریاضی روز به روز در حال گسترش است اما تا کنون استفاده از ریاضیات و آمار فازی در باستان‌شناسی بسیار اندک و انگشت شمار بوده است. از جمله پژوهش‌های انجام شده به موارد زیر اشاره می‌کنیم.

لیفکوفسکیا و همکاران در زمینه‌ی مدل‌های پس‌بینی باستان‌شناسی مطالعاتی انجام دادند. آنها از داده‌های مکانی مربوط به ناحیه‌ای در کشور اسلوواکی استفاده کردند که در آن اعتبارسنجی روش‌های مدل‌سازی پیش‌بینی باستان‌شناسی با استفاده از ریاضیات فازی ارزیابی می‌شود (۱۰). مینک و همکاران مطالعاتی در زمینه مدل‌سازی پیش‌بینی (پس‌بینی) باستان‌شناسی با استفاده از GIS انجام دادند. در این مطالعه ناحیه‌ای در ایالت کنتاکی امریکا انتخاب شده و در آن تحلیل مکانی با مدل‌سازی فازی ترکیب شده و پیش‌بینی‌های مورد نیاز انجام می‌گیرد (۱۱). نیکولوچی و هرمون مطالعه‌ای انجام دادند که در آن گونه‌شناسی باستان‌شناسی با استفاده از

^۳ محمدعلی کبیریان: ma.kabirian@ut.ac.ir

روش‌های فازی، با توجه به چهار مطالعه موردی، بررسی می‌شود. آنها، در مطالعاتی دیگر، درباره‌ی بازسازی مجازی یک بنای فرضی تخریب شده و بازسازی آن با استفاده از قابلیت اطمینان فازی، و فازی‌سازی داده‌های باستان‌شناسی و پیاده‌سازی آن‌ها در GIS که در آن داده‌های استخوانی فازی‌سازی شده و در پایگاه داده‌ها پیاده‌سازی شده‌اند، پژوهش‌هایی انجام داده‌اند و (۱۳) و (۱۲) و (۸). نیز به پژوهش‌دراثر و همکاران اشاره می‌کنیم که در آن با استفاده از منطق فازی داده‌های نادقیق مکانی و زمانی در پایگاه داده‌ها مدیریت می‌شود. در این پژوهش، خیابان‌های یک شهر مربوط به روم باستان با استفاده از نوعی تبدیل فازی در دوره‌ی زمانی مشخصی برآورد می‌شود (۷). از پژوهش‌های داخلی می‌توان به مطالعه‌ی مقصودی و همکاران اشاره کرد که با استفاده از منطق فازی نقش عوامل محیطی را در تعیین مکان‌گزینی سکونتگاه‌های دشت ورامین بررسی کرده‌اند (۵) و پژوهش طاهری و همکاران در زمینه‌ی کاربرد ریاضیات فازی در باستان‌شناسی، که در آن کاربردهای گوناگون ریاضیات فازی در بخش‌های مختلف باستان‌شناسی تشریح شده و نمونه‌ای از این کاربردها بیان شده است (۴).

۲ کاربردهای ریاضیات و آمار فازی در باستان‌شناسی

در وضعیت‌های گوناگونی، استفاده از ریاضیات و آمار فازی در باستان‌شناسی مفید و/یا ضروری است. مهمترین این وضعیت‌ها موارد زیر است.

(الف) هنگامی که داده‌های جمع‌آوری شده/ مستندات/اطلاعات، نادقیق و مبهم هستند.

(ب) هنگامی که روابط بین متغیرها، نادقیق و تقریبی است.

(ج) هنگامی که بین باستان‌شناسان اختلاف نظر وجود دارد.

زمانی که یکی از شرایط بالا فراهم شود می‌توان از ریاضیات فازی در باستان‌شناسی استفاده کرد، اما این استفاده در چه بخش‌هایی از باستان‌شناسی است؟

سه بخش اصلی که می‌توان در آن‌ها از ریاضیات فازی استفاده نمود به شرح زیر است (۴).

(الف) توصیف و تجزیه و تحلیل داده‌ها

(ب) پس‌بینی

(ج) تصمیم‌گیری و نتیجه‌گیری

حال باید سازوکار استفاده از ریاضیات فازی مشخص شود. این سازوکار می‌تواند در قالب قابلیت اطمینان فازی، طراحی سیستم استنتاج فازی و استفاده از آمار

فازی باشد. در ادامه، در بخش ۳ فازی‌سازی داده‌های باستان‌شناسی و در بخش ۴ جدول فراوانی داده‌های فازی‌سازی شده مورد بررسی قرار می‌گیرند.

۳ داده‌های فازی در باستان‌شناسی (مطالعه موردی: داده‌های استخوانی)

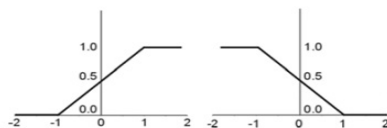
در این بخش داده‌هایی از باستان‌شناسی که قابلیت فازی شدن دارند معرفی شده و سپس چگونگی فازی‌سازی آن‌ها شرح داده شده است. یکی از مهمترین داده‌هایی که باستان‌شناسان با آن سروکار دارند استخوان است. اولین کار باستان‌شناسان هنگام مواجهه با استخوان، تعیین جنسیت، تعیین سن و تاریخگذاری استخوان است. این سه شاخصه یعنی جنسیت، سن و تاریخگذاری در بیشتر مواقع دقیق نیستند و معمولاً دامنه‌ای را در بر می‌گیرند که ابتدا و انتهای آن نیز مبهم است. در ادامه هر یک از این شاخصه‌ها را توضیح می‌دهیم و روند فازی‌سازی آن‌ها را مشخص می‌کنیم.

۱-۳. جنسیت

باستان‌شناسان با استفاده از روش‌هایی مانند مشاهده و بررسی استخوان و یا آزمایش، جنسیت را تعیین می‌کنند اما به علت عوامل گوناگون مانند تخریب استخوان یا کامل نبودن استخوان‌های یک اسکلت در بسیاری از مواقع نمی‌توان دقیقاً جنسیت را تعیین کرد. گاهی مواقع شواهد و نتایج آزمایشگاهی می‌گویند احتمال آنکه استخوان مربوط به یکی از جنس‌های مرد یا زن باشد بیشتر است و گاهی هم بطور کلی نمی‌توان هیچ نظری داد و جنسیت را تعیین کرد. برای آنکه نتایج نادقیق نشوند و محاسبات به نفع یکی از جنسیت‌ها گرد نشوند یا داده‌های نادقیق از چرخه‌ی محاسبات حذف نشوند می‌توانیم آن‌ها را به صورت فازی تعریف کرده و در تحلیل‌ها و نتایج استفاده نماییم.

در ادامه روشی که نیکولوچی و همکاران برای فازی سازی جنسیت داده های استخوانی پیشنهاد داده اند تشریح می شود.

در این روش ضریب جنسیت بین بازه $+2$ تا -2 تعریف شده است. مقادیر مثبت به مرد بودن و مقادیر منفی به زن بودن اشاره می کند یعنی هرچه به $+2$ نزدیک شویم امکان مرد بودن بیشتر است و هرچه به -2 نزدیک شویم امکان زن بودن بیشتر است. مقادیر بیشتر از $+1$ برای مرد بودن و کمتر از -1 برای زن بودن بهترین حالت است که قابلیت اطمینان آن برابر با یک در نظر گرفته می شود. نمودار تابع عضویت زن بودن و تابع عضویت مرد بودن به صورتی که در شکل ۱ نشان داده شده، تعریف شده اند (۱۲).



شکل ۱: تابع عضویت زن بودن ($f(k)$) نمودار سمت راست و تابع عضویت مرد بودن ($m(k)$) نمودار سمت چپ (نیکولوچی و همکاران، ۲۰۰۱)

برای این تابع ضریب استخوان شناسی را k (محور افقی نمودار) تعریف می کنیم که ضریب مرد بودن (m) و ضریب زن بودن (f) با استفاده از آن بدست می آید.

برای نمونه زمانی که مقدار k برابر با 0.5 باشد، ضریب زن بودن (f) برابر 0.25 و ضریب مرد بودن (m) برابر 0.75 است.

در مورد روش مطرح شده باید به نکته ی زیر اشاره نمود.

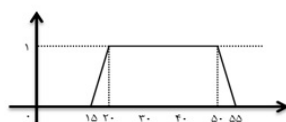
همانطور که پیش از این گفته شد، باستان شناسان برای تشخیص جنسیت از روش مشاهده و بررسی استخوان و آناتومی اسکلت (۳) و همچنین آزمایش، مانند استخراج و بررسی DNA (۶) استفاده می کنند. بدست آوردن ضریب استخوان شناسی (k) بسته به روشی دارد که با استفاده از آن جنسیت را تشخیص می دهیم. برای هر کدام از این روش ها، با توجه به شاخص های گوناگون، سازوکار و تعریف ضریب استخوان شناسی قابل تدوین است که البته نیاز به پژوهشی منحصر به خود دارد. برای نمونه در روش مشاهده استخوان، اندازه ی استخوان و یا استخوان لگن در تشخیص نوع جنسیت موثر است که با توجه به این عوامل می توان سازوکار تعیین ضریب استخوان شناسی را تبیین نمود.

۳-۲. سن

الف) فازی سازی سن بر اساس نظر افراد خبره.

معمولاً در تحلیل هایی که از نتایج آزمایش بر روی استخوان انجام می شود، سن بصورت دقیق بدست نمی آید، بلکه نتیجه ی کار به صورت بازه (فاصله) بیان می شود.

برای نمونه در کاوش های سال ۱۳۸۹ در کورگان های جعفرآباد (کورگان نوعی تدفین به صورت تپه های تدفینی، مربوط به اقوام اورآسیایی است که در شمال غرب ایران یافت می شود) انجام شد، اسکلتی در کورگان شماره ۴ کشف شده بود که بنا به نظر کاوشگر، مربوط به زنی بزرگسال بوده است (۹). برای آنکه سن مشخص شده به همین صورت در تحلیل ها و محاسبات بعدی استفاده شود و مانند روش های سنتی به یک عدد گرد نشود می توان آن را به صورت فازی تعریف کرد. بنا به ایده ی نیکولوچی و همکاران می توان سن استخوان یافت شده را به صورت نمودار دوزنقه ای نشان داد (۱۲). به همین منظور بازه ی سن بزرگسال را 20 تا 50 سال در نظر گرفته (۱۴) و نمودار آن را رسم می کنیم (شکل ۲).



شکل ۲: نمودار تابع عضویت سن فازی حدوداً 20 تا 50 سال

در این نمودار قابلیت اطمینان در بازه ی 20 تا 50 سال یک در نظر گرفته شده است و هرچه از این بازه دور می شویم مقدار قابلیت اطمینان به صفر نزدیک می شود که در نقطه ی 15 و 55 مقدار قابلیت اطمینان صفر می شود.

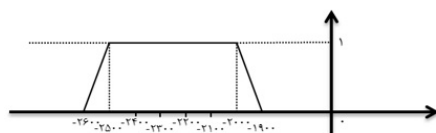
(ب) فازی سازی داده های سن با استفاده از توزیع نرمال.

برای دقیق تر شدن فازی سازی سن می توان از روش های محاسباتی استفاده نمود که نیکولوچی و همکاران آن را در پژوهش خود معرفی کرده اند (۱۲). یکی از این روش ها، روشی است که اسکارسینی و بیگازی به آن پرداخته اند. سازوکار به این صورت است که میانگین (m) و انحراف معیار (s) داده های مورد نظر را بدست می آوریم و با استفاده از آن، تخمین سن، یک توزیع نرمال فرض شود که از آن می توان نتیجه گرفت مقادیر سنی بین m-s و m+s بیشترین امکان را دارند. برای ساده کردن، نمودار را به صورت دوزنقه ای در نظر می گیریم. به طوری که ضریب قابلیت اطمینان در [m-s, m+s] یک خواهد بود و در m-s-q و m+s+q صفر خواهد بود. q یک عدد مثبت است. برای تخمین q، از میان انتخاب های موجود q=0/5s انتخاب می شود. همچنین فرمول عمومی برای f از نظر k و s و m از همین قانون محاسبه می شود که در زیر قابل مشاهده است.

$$f(k) = \begin{cases} 0 & \text{per } k \leq \mu - 1.5\sigma \\ 1 - \frac{(\mu - \sigma) - k}{\sigma/2} & \text{per } \mu - 1.5\sigma < k \leq \mu - \sigma \\ 1 & \text{per } \mu - \sigma < k \leq \mu + \sigma \\ 1 - \frac{k - (\mu + \sigma)}{\sigma/2} & \text{per } \mu + \sigma < k \leq \mu + 1.5\sigma \\ 0 & \text{per } \mu + 1.5\sigma < k \end{cases}$$

۳-۳. تاریخ گذاری

در باستان شناسی روش های گوناگونی مانند استفاده از روش های فیزیکی- شیمیایی، رشد حلقه های درخت (دندروکرونولوژی)، لایه نگاری و غیره برای سالیابی اشیا و داده های بدست آمده وجود دارد که همه ی این روش های علمی و عملی زیر شاخه ی دو روش نظری یعنی سالیابی نسبی و سالیابی مطلق است. سالیابی مطلق اگرچه دقیق تر از سالیابی نسبی است اما نهایتاً در آن می توان یک محدوده ی نسبتاً کوچک برای یک واقعه (مثلاً ساخت سفال یا مرگ یک موجود زنده) نسبت داد. در واقع یک تاریخ مطلق (به معنای واقعی مطلق) شاید تنها به وسیله ی اسناد مکتوب بدست آید (۱). بنابراین باستان شناسان در بیشتر مواقع با تاریخ های نادقیق سروکار دارند. معمولاً تاریخ گذاری دامنه ای زمانی را شامل می شود. برای نمونه نیمه ی دوم هزاره ی سوم پیش از میلاد که به معنی زمانی در فاصله ی [۲۰۰۰- و ۲۵۰۰-] است. بنا به ایده ی نیکولوچی و همکاران می توان هر کدام از بازه های زمانی را به صورت یک مقدار فازی در نمودار دوزنقه ای به شکل زیر نشان داد (شکل ۶) (۱۲).



شکل ۳: تاریخ گذاری فازی

۴. تنظیم جدول فراوانی بر اساس داده های فازی

در این بخش توضیح می دهیم که چگونه می توان جدول فراوانی را بر اساس داده های فازی در مطالعات باستان شناسی تنظیم نمود. در این باره دو نمونه مبتنی بر پژوهش هرمون و همکاران (۸) مطرح می شود.

مثال ۱.۰۴ (داده های سفالین)

جدول ۱ مربوط به داده های حاصل از کاوش در یکی از محوطه های کشور ایتالیا است که در مقاله هرمون و همکاران آمده است. در این جدول نتایج حاصل از دو نوع گونه شناسی فازی و گونه شناسی غیر فازی قابل مشاهده است. گونه شناسی به این صورت است که قطعات سفالی بدست آمده از کاوش با توجه به مشخصاتی که

دارند در گونه‌هایی مانند سبو، کاسه، کوزه و غیره دسته بندی می‌شوند. گاهی این گونه شناسی به راحتی انجام نمی‌گیرد و نمی‌توان به راحتی گونه‌ی قطعه یافت شده را تشخیص داد که در این حالت قطعه یا از آمارگیری کنار گذاشته می‌شود و یا یک گونه برای آن انتخاب شده و به اصطلاح گرد می‌شود. برای جلوگیری از این کار می‌توان گونه‌شناسی فازی را تعریف نمود و داده‌ها را به این صورت فازی سازی کرد و سپس عملیات آماری را روی آن انجام داد. همانطور که مشاهده می‌شود درصدهایی که از گونه‌شناسی غیرفازی و گونه‌شناسی فازی به دست آمده، متفاوت است که این تفاوت در تحلیل‌های حاصل از کاوش مانند اقتصاد معیشتی تغییر ایجاد می‌کند. برای نمونه وجود کوزه‌های ذخیره نشان از ذخیره مواد غذایی دارد که این مساله در جوامع پیش از تاریخ و برای تشخیص ورود انسان از دوران جمع آوری غذا به دوران گردآوری و پس از آن تولید غذا که با دامپروری و کشاورزی همراه است اهمیت دارد.

جدول ۱: درصد فراوانی گونه‌های مختلف سفال‌های نمونه‌گیری شده (۸).

گونه سفال	سبو	کوزه	کاسه	بشقاب	نامشخص
دقیق	۴۴	۳۲	۱۴	۲	۸
فازی	۵۷	۳۴	۸	۱	۰

مثال ۲.۴. (داده‌های استخوان جانوری)

جدول ۲ مربوط به داده‌های حاصل از کاوش ویلایی روستایی در بریقه لیبی است که در مقاله هرمون و همکاران آمده است. داده‌ها، استخوان‌های جانوری هستند که با توجه سن گونه‌ی جانوری دسته‌بندی شده‌اند. برای هر کدام از گروه‌های سنی از قبل مقداری تعریف شده و داده‌ها با توجه به آنها دسته بندی می‌شوند. اما همیشه سن یافته‌ها دقیق بدست نمی‌آید و برای آنکه همه‌ی یافته‌ها در تحلیل جای بگیرند از روش‌های فازی و فازی سازی داده‌ها استفاده کرده و سپس عملیات آماری بر روی داده‌ها انجام می‌شود. همانطور که مشاهده می‌شود نتایجی که از دو روش فازی و دقیق بدست آمده، متفاوت است که این تفاوت باعث تغییر در تحلیل‌ها می‌شود. یکی از شاخص‌هایی که می‌توان با آن اقتصاد معیشتی جوامع را بررسی نمود نوع خوراک آن جوامع است. جوامع انسانی از گذشته تا حال از دام استفاده‌های گوناگونی داشته‌اند. از جمله استفاده از گوشت، شیر، پشم، نیروی محرکه‌ی دام برای کشاورزی است. فراوانی دام در سنین مختلف گویای مطالب گوناگونی است. در این مطالعه همانطور که مشاهده می‌شود هنگام استفاده از روش سنتی و دقیق، داده‌ی استخوانی در سن آغاز بزرگسالی وجود ندارد که نشان می‌دهد بطور معمول در این سن دام‌ها کشته نمی‌شدند و فرآورده‌ای مانند شیر در این جامعه هدف اصلی از پرورش دام بوده است اما هنگامی که با روش فازی داده‌ها تحلیل می‌شود، حدود ۲۰ درصد داده‌ها در بازه سنی آغاز بزرگسالی قرار می‌گیرد که درصد قابل توجهی است و نشان دهنده‌ی آن است که تامین گوشت از طریق دامپروری مورد توجه بوده است.

جدول ۲: درصد فراوانی گونه‌های مختلف سفال‌های نمونه‌گیری شد (۸).

گروه سنی	بسیار جوان (۰ تا ۶ ماه)	جوان (۶ تا ۲۴ ماه)	آغاز بزرگسالی (۲ تا ۴ سال)	بزرگسال (۴ تا ۸ سال)	کهنسال (بزرگتر از ۸ سال)
دقیق	۸	۳۰	۰	۵۴	۸
فازی	۵	۳۹	۲۰	۳۰	۶

نتیجه‌گیری

وجود داده‌های نادقیق در باستان‌شناسی زمینه‌ی استفاده از ریاضیات فازی در باستان‌شناسی را فراهم می‌کند. از مواردی که می‌توان با استفاده از آن به باستان‌شناسان و تحلیل‌هایشان کمک نمود استفاده از آمار فازی است. به منظور استفاده از آمار فازی در باستان‌شناسی در آغاز باید داده‌های باستان‌شناسی فازی سازی شود. از جمله داده‌هایی که می‌توان این فرآیند را بر روی آن‌ها انجام داد، سن، جنسیت و تاریخگذاری استخوان است. با فازی‌سازی این داده‌ها، زمینه‌ی استفاده از آن‌ها در آمار فازی

فراهم می‌شود. نتایجی که با استفاده از آمار فازی از داده‌های باستان‌شناسی بدست می‌آید با نتایجی که با استفاده از روش‌های غیر فازی بدست آمده متفاوت است و همین تفاوت باعث تغییر در تحلیل‌های باستان‌شناسان می‌شود که اهمیت استفاده از روش‌های فازی در باستان‌شناسی را نشان می‌دهد.

مراجع

[۱] بحرالعلومی شاپورآبادی، ف. (۱۳۹۲). *روشهای سالیابی در باستان‌شناسی*، انتشارات سمت، تهران.

[۲] دارک، کن.آر. ترجمه‌ی عبدی، ک. (۱۳۹۴). *مبانی نظری باستان‌شناسی*، مرکز نشر دانشگاهی، تهران.

[۳] می‌ز، سایمون. ترجمه‌ی اشرفیان بناب، م. (۱۳۸۱). *باستان‌شناسی/استخوان‌های انسان*، انتشارات سازمان میراث فرهنگی کشور، تهران.

[۴] طاهری، م. و ایروانی قدیم، ف. و کبیریان، م. (۱۳۹۶). استفاده از ریاضیات فازی در باستان‌شناسی، ششمین کنگره مشترک سیستم‌های فازی و هوشمند ایران، کرمان، ۳۰۷-۳۱۴.

[۵] مقصودی، م. و همکاران. (۱۳۹۴). تحلیل نقش عوامل محیطی در مکانگزینی سکونتگاههای پیش از تاریخ دشت ورامین با استفاده از منطق فازی، برنامه ریزی و آمایش فضا، ۱۹ (۳)، ۲۶۱-۲۳۳.

[۶] نیکنامی، ک. و رمضانی، م. (۱۳۹۵). *مقدمه‌ای بر ژنتیک باستان‌شناسی*، انتشارات سمت، تهران.

[7] De Runz C, and Desjardin E, and Piantoni F. and Herbin M. (2006), Using fuzzy logic to manage uncertain multi-modal data in an archaeological GIS, *International Symposium on Spatial Data Quality-ISSDQ*. Volume 7, 1-4.

[8] Hermon S, and Niccolucci F, and Alhaique F, and Iovini M, and Leonini V. (2004), Archaeological typologies-an archaeological fuzzy reality, . *British Archaeological Reports Internatianl Series* .1227, 30-34.

[9] Iravani Ghadim F. (2014), Jafar Abad Kurgan no IV, *Veli Sevin'e Armağan*, 87-106.

[10] Lieskovská T, and Ďuračiová R, and Karella L. (2013), Selected Mathematical Principles of Archaeological Predictive Models, *Interdisciplinaria Archaeologica*. 5 , 177-190 .

[11] Mink P, and Ripy J, and Bailey K, and Grossardt T. (2009), Predictive Archaeological Modeling Using GIS-Based Fuzzy Set Estimation: Case Study of Woodford County, *Kentucky Transportation Center Faculty and Researcher Publications*. paper 12 .

[12] Niccolucci F, and Andrea A.D, and Crescioli M. (2011), Archaeological applications of fuzzy databases, *Archaeological Reports Internatianl Series*. 931, 107-116.

[13] Niccolucci F, and Hermon S. (2004), A fuzzy logic approach to reliability in archaeological virtual reconstruction., the Artifact-Digital Interpretation of the Past, (Proceedings of the CAA2004 Conference). 26-33.

[14] White T.D, and Black M.T, and Folkens P.A. (2014), *Human Osteology*, 3rd Edition, Academic Press.

رگرسیون خطی ساده چندکی فازی

محسن عارفی^۱

گروه آمار، دانشکده علوم ریاضی و آمار، دانشگاه بیرجند

چکیده

یک شیوه جدید برای برازش مدل رگرسیونی خطی ساده چندکی بر اساس متغیر مستقل معمولی، متغیر پاسخ فازی و ضرایب فازی معرفی شده است. در ابتدا، یک تابع زیان بین تفاضل دو عدد فازی معرفی می‌گردد و آنگاه بر اساس آن، ضرایب مدل رگرسیونی با استفاده از روش کمترین چندک خطا به داده‌های فازی برازش می‌گردد. قابل توجه است که روش پیشنهادی تحت داده‌های پرت استوار است. در نهایت، روش پیشنهادی، بر روی یک مثال عددی با وجود داده‌های پرت مختلف مورد بررسی قرار گرفته است.

کلمات کلیدی: داده فازی، رگرسیون خطی ساده، کمترین چندک خطا، تابع زیان.

۱ مقدمه

یکی از مسائل مهم در مباحث آماری، مسئله برازش یک مدل رگرسیونی بین متغیر پاسخ و برخی متغیرهای مستقل است. با این وجود در برخی موارد، برخی کمیت‌های موجود به صورت نادقیق (فازی) گزارش می‌شوند و ما باید در چنین شرایطی یک مدل رگرسیونی مناسب به داده‌ها برازش دهیم. از طرف دیگر، ممکن است در بین داده‌های موجود، داده‌های پرتی نیز مشاهده گردند. در چنین مواردی، یک گزینه مناسب استفاده از مدل‌های رگرسیونی استوار در محیط فازی است. این مدل‌های استوار می‌تواند بر اساس روش‌های کمترین چندک خطای بین متغیر پاسخ مشاهده شده و متغیر پاسخ برآورد شده بنا گردند. در سال‌های اخیر، نویسندگان متعددی به بررسی و تحقیق مبحث رگرسیون استوار در محیط فازی پرداخته‌اند. در زیر به برخی از آنها اشاره شده است.

سون (۲۰۰۵) بر اساس M-برآوردگرها، برخی مدل‌های رگرسیون استوار را به دست آورده است. وارگا (۲۰۰۷) دو برآورد استوار از ضرایب فازی در مدل رگرسیونی ارائه داده است. نصرآبادی و همکاران (۲۰۰۷) یک نگرش برای تعیین داده‌های پرت در مدل‌های رگرسیونی فازی پیشنهاد نموده‌اند. کولا و آپیدین (۲۰۰۸) بر اساس رتبه‌بندی بین مجموعه‌های فازی به ارائه برخی مدل‌های رگرسیونی استوار پرداخته‌اند. برخی مدل‌های رگرسیونی استوار بر اساس برآورد کمترین میانه مربعات وزنی توسط دورسو و همکاران (۲۰۱۱) مورد تحقیق قرار گرفته است. مدل‌های C-رگرسیون فازی استوار به وسیله لسکی و کوتای (۲۰۱۵) مورد مطالعه قرار گرفته است. برای

^۱محسن عارفی: Arefi@Birjand.ac.ir

بررسی برخی نگرشهای دیگر در زمینه مدل‌های رگرسیونی فازی استوار به چاچی و طاهری (۲۰۱۳)، یانگ و همکاران (۲۰۱۳)، چاچی و همکاران (۲۰۱۴)، زانگ و لو (۲۰۱۶)، چاچی و روزبه (۲۰۱۷) مراجعه نمایید.

در این مقاله، یک شیوه جدید برای برآورد ضرایب فازی یک مدل رگرسیونی خطی ساده چندکی بر اساس یک تابع زیان بین تفاضل دو عدد فازی، ارائه شده است. مزیت مدل پیشنهادی در این است که تحت داده‌های پرت استوار است. در ادامه، نیز یک ملاک برای نیکویی برازش مدل معرفی و مدل مورد ارزیابی قرار گرفته است. در زیر به معرفی برخی مفاهیم اساسی مورد استفاده در مقاله پرداخته شده است (طاهری؛ ۱۳۷۵).

یک مجموعه فازی $\tilde{A}$ روی مجموعه مرجع X با تابع عضویت $\tilde{A} : X \rightarrow [0, 1]$ تعریف می‌گردد. این مجموعه فازی را یک عدد فازی مثلثی گویند، اگر تابع عضویت آن به صورت زیر باشد:

$$\tilde{A}(x) = \begin{cases} 1 + \frac{x-m}{r} & m-r \leq x \leq m, \\ 1 - \frac{x-m}{s} & m < x \leq m+s. \end{cases}$$

یک عدد فازی مثلثی را به صورت $\tilde{A} = (m, r, s)_T$ نشان می‌دهند. اگر $r = s$ باشد، آنرا عدد فازی مثلثی متقارن گویند و به صورت $\tilde{A} = (m, r)_T$ نشان می‌دهند. عملگرهای جمع و ضرب بین دو عدد فازی $\tilde{M}_1 = (m_1, r_1, s_1)_T$ و $\tilde{M}_2 = (m_2, r_2, s_2)_T$ به صورت زیر تعریف می‌شود:

$$\begin{aligned} \widetilde{M}_1 \oplus \widetilde{M}_2 &= (m_1 + m_2, r_1 + r_2, s_1 + s_2)_T, \\ \lambda \otimes \widetilde{M}_1 &= \begin{cases} (\lambda m_1, \lambda r_1, \lambda s_1)_T & \lambda > 0, \\ (\lambda m_1, -\lambda s_1, -\lambda r_1)_T & \lambda < 0. \end{cases} \end{aligned}$$

۲ رگرسیون خطی ساده چندکی فازی

در این بخش، می‌خواهیم پارامترهای یک مدل رگرسیونی ساده را طوری برآورد نماییم که تحت داده‌های پرت استوار باشد. برای این کار، بجای استفاده از روش حداقل میانگین خطا، از روش حداقل چندک خطا استفاده می‌شود. در ابتدا نیاز به معرفی یک تابع زیان بین تفاضل دو عدد فازی مثلثی به صورت زیر هستیم.

تعریف ۱.۲. فرض کنید $\tilde{M}_1 = (m_1, \alpha_1)_T$ و $\tilde{M}_2 = (m_2, \alpha_2)_T$ دو عدد فازی مثلثی متقارن باشند. آنگاه تابع زیان تفاضل این دو عدد فازی مثلثی به صورت زیر تعریف می‌گردد:

$$\rho_\tau(\tilde{M}_1 - \tilde{M}_2) = \frac{1}{3} \left[\rho_\tau(m_1 - m_2) + \rho_\tau((m_1 - m_2) - \frac{1}{\tau}(\alpha_1 - \alpha_2)) + \rho_\tau((m_1 - m_2) + \frac{1}{\tau}(\alpha_1 - \alpha_2)) \right],$$

که در آن $\tau \in (0, 1)$ ، $\rho_\tau(u) = u(\tau - I(u < 0))$ و $I(\cdot)$ تابع نشانگر است.

فرض کنید که یک مجموعه داده به صورت $(x_i, \tilde{Y}_i)$ ، $i = 1, 2, \dots, n$ در دسترس باشد، که در آن متغیر پاسخ یک عدد فازی مثلثی متقارن به صورت $\tilde{Y}_i = (y_i, r_i)_T$ می‌باشد. اکنون ما می‌خواهیم یک مدل رگرسیونی فازی به صورت زیر به داده‌ها برازش دهیم:

$$\hat{\tilde{Y}} = \tilde{A}_0 \oplus (\tilde{A}_1 \otimes x),$$

که در آن ضرایب به صورت اعداد فازی مثلثی متقارن $\tilde{A}_j = (a_j, \alpha_j)_T$ ، $j = 0, 1$ هستند. بر اساس عملگرهای حسابی بین اعداد فازی، متغیر پاسخ به صورت زیر برآورد می‌شود (فرض کنید $x_i > 0$):

$$\begin{aligned} \hat{\tilde{Y}}_i &= \tilde{A}_0 \oplus (\tilde{A}_1 \otimes x) \\ &= (a_0 + a_1 x_i; \alpha_0 + \alpha_1 x_i)_T \\ &= (\hat{y}_i; \hat{r}_i)_T, \quad i = 1, 2, \dots, n, \end{aligned}$$

مجموع تابع زیان تفاضل بین متغیر پاسخ مشاهده شده و متغیر پاسخ برآورد شده به صورت زیر معرفی می‌گردد:

$$\begin{aligned} H &= \sum_{i=1}^n \rho_{\tau}(\tilde{Y}_i - \hat{Y}_i) \\ &= \frac{1}{3} \sum_{i=1}^n \left[\rho_{\tau}(y_i - \hat{y}_i) + \rho_{\tau}((y_i - \hat{y}_i) - \frac{1}{3}(r_i - \hat{r}_i)) + \rho_{\tau}((y_i - \hat{y}_i) + \frac{1}{3}(r_i - \hat{r}_i)) \right]. \end{aligned}$$

با مینیم کردن تابع H ، برآورد ضرایب مدل رگرسیونی به دست می‌آید. برای مینیم کردن تابع H ، آن را به یک مسئله برنامه‌ریزی خطی تبدیل می‌کنیم. فرض کنید $w_i = (y_i - \hat{y}_i) + l(r_i - \hat{r}_i)$ و $v_i = (y_i - \hat{y}_i) - l(r_i - \hat{r}_i)$ ، $e_i = y_i - \hat{y}_i$ زیر بازنویسی نمائیم:

$$\begin{aligned} H &= \frac{1}{3} \left[\sum_{i=1}^n \rho_{\tau}(e_i) + \sum_{i=1}^n \rho_{\tau}(v_i) + \sum_{i=1}^n \rho_{\tau}(w_i) \right] \\ &= \frac{1}{3} \left[\sum_{i=1}^n [e_i(\tau - I(e_i < 0))] + \sum_{i=1}^n [v_i(\tau - I(v_i < 0))] + \sum_{i=1}^n [w_i(\tau - I(w_i < 0))] \right] \\ &= \frac{1}{3} \left[\sum_{i=1}^n [e_i^-(1 - \tau) + e_i^+ \tau] + \sum_{i=1}^n [v_i^-(1 - \tau) + v_i^+ \tau] + \sum_{i=1}^n [w_i^-(1 - \tau) + w_i^+ \tau] \right] \\ &= \frac{1}{3} \left[\sum_{i=1}^n [(e_i^- + v_i^- + w_i^-)(1 - \tau) + (e_i^+ + v_i^+ + w_i^+) \tau] \right] \end{aligned}$$

که در آن داریم:

$$\begin{aligned} e_i^- &= -\min(e_i, 0), & e_i^+ &= \max(e_i, 0), & e_i &= e_i^+ - e_i^-, \\ v_i^- &= -\min(v_i, 0), & v_i^+ &= \max(v_i, 0), & v_i &= v_i^+ - v_i^-, \\ w_i^- &= -\min(w_i, 0), & w_i^+ &= \max(w_i, 0), & w_i &= w_i^+ - w_i^-. \end{aligned}$$

در نتیجه، مسئله مینیم کردن تابع H ، با حل مسئله برنامه‌ریزی خطی زیر میسر می‌گردد:

$$\begin{aligned} \min \quad H &= \frac{1}{3} \left[\sum_{i=1}^n [(e_i^- + v_i^- + w_i^-)(1 - \tau) + (e_i^+ + v_i^+ + w_i^+) \tau] \right] \\ \text{s.t.} \end{aligned}$$

$$\begin{aligned} y_i - a_0 - a_1 x_i &= e_i^+ - e_i^-, \\ (y_i - a_0 - a_1 x_i) - \frac{1}{3}(r_i - \alpha_0 - \alpha_1 x_i) &= v_i^+ - v_i^-, \\ (y_i - a_0 - a_1 x_i) + \frac{1}{3}(r_i - \alpha_0 - \alpha_1 x_i) &= w_i^+ - w_i^-, \\ r_i - \hat{r}_i &\geq 0, \quad i = 1, \dots, n, \\ e_i^-, e_i^+, v_i^-, v_i^+, w_i^-, w_i^+ &\geq 0, \quad i = 1, \dots, n, \\ a_j \in R, \alpha_j &\geq 0, \quad j = 0, 1. \end{aligned}$$

تذکر ۲.۰۲. قابل توجه است که مدل‌های رگرسیونی چندکی، از حداقل چندک خطا (تفاضل بین متغیر پاسخ مشاهده شده و برآورد شده) به جای میانگین (مجموع) حداقل مربعات خطا استفاده می‌کنند. در این مدل‌ها، تابع $\rho_{\tau}(\cdot)$ نقش اساسی را در بدست آوردن چندک τ ایفا می‌نماید. از طرف دیگر، استفاده از تابع $\rho(\cdot)$ کار محاسبه ضرایب مدل را ساده‌تر می‌کند (برای بررسی بیشتر، به فصل ۱ کوئنکر (۲۰۰۵) مراجعه نمائید). از آنجایی که چندک (به خصوص میانه و چارکهای اول و سوم) نسبت به میانگین تحت تأثیر داده پرت قرار نمی‌گیرد، بنابراین، روش بیان شده نسبت به وجود داده پرت حساس نیست و جوابهای بهینه‌ای را می‌دهد.

تذکر ۳.۰۲. برای حل مسئله برنامه‌ریزی خطی فوق، از نرم افزار *LINGO* استفاده شده است.

۱.۲ نیکویی برازش مدل رگرسیونی

برای بررسی ارزیابی نیکویی برازش مدل‌های رگرسیونی، از میانگین تابع زیان تفاضل بین مقادیر متغیر پاسخ مشاهده شده و برآورد شده به صورت زیر استفاده نمود. هر چه این مقدار کوچکتر باشد، برازش مدل‌ها به داده‌های موجود بهتر است.

تعریف ۴.۲. فرض کنید $\tilde{Y}_i$ و $\hat{\tilde{Y}}_i$ به ترتیب مقادیر متغیر پاسخ مشاهده شده و مقادیر متغیر پاسخ برآورد شده باشند. میانگین تابع زیان تفاضل بین مقادیر $\tilde{Y}_i$ و $\hat{\tilde{Y}}_i$ به صورت زیر تعریف می‌شود:

$$M = \frac{1}{n} \sum_{i=1}^n \rho_{\tau}(\tilde{Y}_i - \hat{\tilde{Y}}_i).$$

۳ مثال عددی

داده‌های موجود در جدول ۱ را در نظر بگیرید. فرض کنید بخواهیم یک مدل رگرسیونی چندکی بین متغیر پاسخ فاز $\tilde{Y}_i$ و متغیر مستقل x_i به صورت زیر برازش دهیم:

$$\hat{\tilde{Y}}_i = \tilde{A}_0 \oplus (\tilde{A}_1 \otimes x_i), \quad i = 1, \dots, 10$$

که در آن ضرایب $\tilde{A}_0 = (a_0, \alpha_0)_T$ و $\tilde{A}_1 = (a_1, \alpha_1)_T$ به صورت اعداد فاز $\tilde{A}_1$ و $\tilde{A}_0$ متناهی مقارن هستند.

بر اساس بخش قبل، برای بدست آوردن ضرایب مدل رگرسیونی، باید مسئله برنامه ریزی خطی زیر را حل نماییم:

جدول ۱: مجموعه داده‌های مشاهده شده

$\tilde{Y}_i = (y_i, r_i, s_i)_T$	x_i	$\tilde{Y}_i = (y_i, r_i, s_i)_T$	x_i
$(6, 0/6)_T$	6	$(2, 0/2)_T$	1
$(6, 0/6)_T$	7	$(3, 0/3)_T$	2
$(7, 0/7)_T$	8	$(5, 0/5)_T$	3
$(8, 0/8)_T$	9	$(4, 0/4)_T$	4
$(10, 1/0)_T$	10	$(5/5, 0/6)_T$	5

$$\min H = \frac{1}{\tau} [\sum_{i=1}^n [(e_i^- + v_i^- + w_i^-)(1 - \tau) + (e_i^+ + v_i^+ + w_i^+)\tau]]$$

s.t.

$$2 - a_0 - a_1 = e_1^+ - e_1^-,$$

$$2 - a_0 - a_1 - \frac{1}{\tau}(0/2 - \alpha_0 - \alpha_1) = v_1^+ - v_1^-,$$

$$2 - a_0 - a_1 + \frac{1}{\tau}(0/2 - \alpha_0 - \alpha_1) = w_1^+ - w_1^-,$$

$$0/2 - \alpha_0 - \alpha_1 \geq 0,$$

:

$$10 - a_0 - 10a_1 = e_{10}^+ - e_{10}^-,$$

$$10 - a_0 - 10a_1 - \frac{1}{\tau}(1 - \alpha_0 - 10\alpha_1) = v_{10}^+ - v_{10}^-,$$

$$10 - a_0 - 10a_1 + \frac{1}{\tau}(1 - \alpha_0 - 10\alpha_1) = w_{10}^+ - w_{10}^-,$$

$$e_i^-, e_i^+, v_i^-, v_i^+, w_i^-, w_i^+ \geq 0, \quad i = 1, \dots, 10,$$

$$1 - \alpha_0 - 10\alpha_1 \geq 0,$$

$$a_0 \in R, a_1 \in R, \alpha_0 \geq 0, \alpha_1 \geq 0.$$

۰/۷۵, ۰/۵, ۰/۲۵ $\tau =$ در جدول ۲ خلاصه شده است. با توجه به ملاک نیکویی برازش M ، مدل رگرسیونی چندکی با $\tau = ۰/۲۵$ به داده‌ها برازش بهتری داشته است (شکل ۱). اکنون فرض کنید تغییراتی را در داده‌ها ایجاد نماییم و استوار بودن مدل رگرسیون چندکی را بررسی کنیم.

جدول ۲: مدل‌های رگرسیونی به ازای چندکهای مختلف در داده‌های اصلی

چندک	مدل رگرسیونی	M
$\tau = ۰/۲۵$	$(۰/۷۱۴۳, ۰/۰۷۱۴)_T \otimes x \oplus (۱/۲۸۵۷, ۰/۱۲۸۶)_T$	۰/۱۵۱۸
$\tau = ۰/۵۰$	$(۰/۷۱۴۳, ۰/۰۷۱۴)_T \otimes x \oplus (۱/۵۷۱۴, ۰/۱۵۷۱)_T$	۰/۲۳۲۱
$\tau = ۰/۷۵$	$(۰/۸۷۵۰, ۰/۰۸۷۵)_T \otimes x \oplus (۱/۲۵۰۰, ۰/۱۲۵۰)_T$	۰/۲۱۵۶

الف) وجود داده پرت: فرض کنید در مشاهده پنجم، متغیر پاسخ به صورت داده پرت $(۱۵, ۰/۶)_T$ تغییر کند (یعنی مرکز متغیر پاسخ بجای مقدار ۵/۵ به مقدار ۱۵ تغییر کرده است). انواع مدل‌های رگرسیونی به ازای $\tau = ۰/۷۵, ۰/۵, ۰/۲۵$ در جدول ۳ خلاصه شده است. با توجه به ملاک نیکویی برازش M ، مدل رگرسیونی چندکی با $\tau = ۰/۲۵$ به داده‌ها برازش بهتری داشته است و از طرفی مدل رگرسیونی آن نیز مشابه با مدل رگرسیونی مبتنی بر داده‌های اصلی است. یعنی مدل تحت داده پرت قرار نگرفته است (شکل ۱).

جدول ۳: مدل‌های رگرسیونی به ازای چندکهای مختلف با وجود داده پرت در قسمت الف

چندک	مدل رگرسیونی	M
$\tau = ۰/۲۵$	$(۰/۷۱۴۳, ۰/۰۷۱۴)_T \otimes x \oplus (۱/۲۸۵۷, ۰/۱۲۸۶)_T$	۰/۳۸۹۳
$\tau = ۰/۵۰$	$(۰/۷۱۴۳, ۰/۰۷۱۴)_T \otimes x \oplus (۱/۵۷۱۴, ۰/۱۵۷۱)_T$	۰/۷۰۷۱
$\tau = ۰/۷۵$	$(۰/۷۱۴۳, ۰/۰۷۱۴)_T \otimes x \oplus (۲/۸۵۷۱, ۰/۲۸۵۷)_T$	۰/۹۰۳۶

ب) افزایش پهنای یک داده: فرض کنید در مشاهده پنجم، متغیر پاسخ به صورت $(۵/۵, ۵)_T$ تغییر کند (یعنی پهنای متغیر پاسخ بجای مقدار ۰/۶ به مقدار ۵ تغییر کرده است). انواع مدل‌های رگرسیونی به ازای $\tau = ۰/۷۵, ۰/۵, ۰/۲۵$ در جدول ۴ خلاصه شده است. با توجه به ملاک نیکویی برازش M ، مدل رگرسیونی چندکی با $\tau = ۰/۲۵$ به داده‌ها برازش بهتری داشته است. با وجود تغییر زیاد پهنای متغیر پاسخ پنجمین مشاهده، مدل رگرسیونی تغییر زیادی از خود نشان نداده و استوار بودن خود را حفظ می‌کند (شکل ۱).

جدول ۴: مدل‌های رگرسیونی به ازای چندکهای مختلف در قسمت ب

چندک	مدل رگرسیونی	M
$\tau = ۰/۲۵$	$(۰/۷۳۱۳, ۰/۰۳۷۵)_T \otimes x \oplus (۱/۱۵۰۰, ۰/۴۰۰۰)_T$	۰/۲۰۴۳
$\tau = ۰/۵۰$	$(۰/۷۳۱۳, ۰/۰۳۷۵)_T \otimes x \oplus (۱/۴۱۸۷, ۰/۴۶۲۵)_T$	۰/۲۹۳۰
$\tau = ۰/۷۵$	$(۰/۸۷۵۰, ۰/۰۸۷۵)_T \otimes x \oplus (۱/۲۵۰۰, ۰/۱۲۵۰)_T$	۰/۲۸۵۴

ج) وجود دو داده پرت در دو جهت مختلف: فرض کنید متغیر پاسخ در مشاهده پنجم به صورت $(۱۵, ۰/۶)_T$ و متغیر پاسخ در مشاهده هشتم به صورت $(۰, ۰/۷)_T$ تغییر کنند. انواع مدل‌های رگرسیونی به ازای $\tau = ۰/۷۵, ۰/۵, ۰/۲۵$ در جدول ۵ خلاصه شده است. با توجه به ملاک نیکویی برازش M ، مدل رگرسیونی چندکی با $\tau = ۰/۲۵$ به داده‌ها برازش بهتری داشته است. با اینکه دو داده پرت در بین داده وجود دارد، ولی مدل رگرسیونی تغییر زیادی از خود نشان نداده و استوار مانده است (شکل ۱).

مراجع

[۱] طاهری، س.م. (۱۳۷۵)، آشنایی با نظریه مجموعه‌های فازی. انتشارات جهاد دانشگاهی، دانشگاه فردوسی مشهد.

جدول ۵: مدلهای رگرسیونی به ازای چندکهای مختلف در قسمت ج

M	مدل رگرسیونی	چندک
$\circ/۸۹۱۷$	$(۱/۳۳۳, \circ/۱۳۳۳)_T \oplus (\circ/۶۶۶۷, \circ/۶۶۶۷)_T \otimes x$	$\tau = \circ/۲۵$
$۱/۰۵۷۱$	$(۱/۵۷۱۴, \circ/۱۵۷۱)_T \oplus (\circ/۷۱۴۳, \circ/۰۷۱۴)_T \otimes x$	$\tau = \circ/۵۰$
$۱/۰۷۸۶$	$(۲/۸۵۷۱, \circ/۲۸۵۷)_T \oplus (\circ/۷۱۴۳, \circ/۰۷۱۴)_T \otimes x$	$\tau = \circ/۷۵$

شکل ۱: مدلهای رگرسیونی چندکی در حالت‌های مختلف با $\tau = \circ/۲۵$

- [2] Chachi, J. and Roozbeh, M. (2017), A fuzzy robust regression approach applied to bedload transport data, *Communications in Statistics - Simulation and Computation*, **46**, 1703-1714.
- [3] Chachi, J. and Taheri, S.M. (2013), A least-absolutes regression model for imprecise response based on the generalized Hausdorff-metric, *Journal of Uncertain Systems*, **7**, 265-276.
- [4] Chachi, J., Taheri, S.M., and Arghami, N.R. (2014), A hybrid fuzzy regression model and its application in hydrology engineering, *Applied Soft Computing*, **25**, 149-158.
- [5] D'Urso, P., Massari, R., and Santoro, A. (2011), Robust fuzzy regression analysis. *Information Sciences*, **181**, 4154-4174.
- [6] Koenker, R. (2005), *Quantile Regression*, Cambridge, New York.
- [7] Kula, K.S. and Apaydin, A. (2008), Fuzzy robust regression analysis based on the ranking of fuzzy sets. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, **16**, 663-681.
- [8] Leski, J.M. and Kotas M. (2015), On robust fuzzy c-regression models, *Fuzzy Sets and Systems*, **279**, 112-129.
- [9] Nasrabadi, E., Hashemi, S.M., and Ghatee, M. (2007), An LP-based approach to outliers detection in fuzzy regression analysis. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, **15**, 441-456.
- [10] Sohn, B. (2005), Robust fuzzy linear regression based on M-estimators, *Journal of Applied Mathematics and Computing*, **18**, 591-601.
- [11] Varga, S. (2007), Robust estimations in classical regression models versus robust estimations in fuzzy regression models, *Kybernetika*, **43**, 503-508.
- [12] Yang, Z., Yin, Y., and Chen, Y. (2013), Robust fuzzy varying coefficient regression analysis with crisp inputs and Gaussian fuzzy output, *Journal of Computing Science and Engineering*, **7**, 263-271.
- [13] Zhang, D. and Lu, Q. (2016), Robust regression analysis with LR-type fuzzy input variables and fuzzy output variable, *Journal of Data Analysis and Information Processing*, **4**, 64-80.

دسترس پذیری سیستم های مهندسی با طول عمر فازی

زهره عباس نیا^۱، خدیجه توکلی ناربند^۲، بهرام صادقیپور گیلده^۳

^{۱،۲} کارشناس ارشد آمار ریاضی دانشگاه فردوسی مشهد

^۳ استاد گروه آمار دانشگاه فردوسی مشهد

چکیده

قابلیت اطمینان را می توان به طور کلی به عنوان میزان توانایی یک سیستم در انجام اهداف خواسته شده از آن در مدت زمان معینی تعریف نمود. یکی از شاخص های مهم قابلیت اطمینان، دسترس پذیری سیستم ها است. مفهوم دسترس پذیری، در اصل برای سیستم های تعمیرپذیر که نیازمند کار مداوم هستند، توسعه یافته است. دسترس پذیری، احتمال این که یک سیستم در نقطه زمانی t به صورت مطلوب کار کند، تعریف می شود. در سیستم های واقعی همواره با مواردی مواجه هستیم که داده ها به صورت نادقیق هستند. در چنین موقعیت هایی برای تصمیم گیری در مورد وضعیت سیستم، تلفیق روش های معمول با نظریه مجموعه های فازی پیشنهاد می شود. در این مقاله با ارزیابی سیستم های با طول عمر فازی کمک می شود تا با بهره گیری از علم فازی، تصمیمات صحیح و بهنگام در مورد واحدهای عملیاتی متشکل از سیستم های پیچیده اتخاذ شود. نتایج عددی نشان داد که در مقایسه با روش های معمول بازه های دقیق تری برای قابلیت اطمینان و دسترس پذیری سیستم های مهندسی به دست می آید.

کلمات کلیدی: اعداد فازی، دسترس پذیری، سیستم های تعمیرپذیر، سیستم های مهندسی و قابلیت اطمینان.

۱ مقدمه

با پیشرفت روز افزون تکنولوژی، بهره گیری از سیستم های پیچیده که از اجزای گوناگونی تشکیل شده است، امری اجتناب ناپذیر است؛ بررسی اجزایی که دارای عملکردهای متفاوتی هستند، به ویژه وقتی که طول عمر آن ها دقیقاً مشخص نیست، در قالب داده های فازی بیان می شوند. از مباحث مهم قابلیت اطمینان سامانه ها، دسترس پذیری یک سیستم یا واحد که ارتباط مستقیمی با طول عمر آن دارد، از جمله موضوعاتی است که امروزه دغدغه فکری بسیاری از متخصصین صنایع مهم و پیشرفته شده است، زیرا که اطلاع از دسترس پذیری سیستم ها و تحقیق در زمینه ی طول عمر قطعات منجر به این خواهد شد که کارشناسان و متخصصین قبل از خرابی قطعات و ازکارافتادگی کل سیستم، قطعات فرسوده را با قطعات سالم جایگزین نموده و به این ترتیب از ایجاد خلل در چرخه ی عرضه و تقاضا و تحمیل هزینه های گزاف به صنایع جلوگیری نمایند.

امروزه تحقیق در زمینه ی ارزیابی عملکرد سیستم ها و محاسبه ی دسترس پذیری سیستم های مهندسی در تمام جنبه های کاربردی به صورت یک اصل و نیاز اساسی مورد پذیرش محققین، پژوهشگران و مدیران بهره بردار از صنایع گوناگون می باشد. بررسی های انجام شده حاکی از آن است که برای طول عمر واحدها عملاً نمی توان مقدار دقیق و مشخصی در نظر گرفت و در واقع

^۱ ارائه دهنده مقاله: زهره عباس نیا، پست الکترونیک: abbasnia.zohre@mail.um.ac.ir
^۲ نویسنده مسئول: بهرام صادقیپور گیلده، پست الکترونیک: sadeghpour@um.ac.ir

طول عمر واحدها اعداد نامشخص و نادقیق (فازی) می‌باشد. لذا می‌بایست برای قابلیت اطمینان و دسترس‌پذیری آن‌ها چاره‌ای اندیشیده شود. یکی از مهمترین موضوعات مرتبط با موارد فوق، محاسبه‌ی مدت زمان عملیاتی بودن یک سیستم و پیش‌بینی زمان خرابی و ازکارافتادگی سیستم می‌باشد، با این حال که طول عمر قطعات نامشخص است. در رابطه با دسترس‌پذیری سیستم‌های مهندسی افرادی چون کوی و لی (۲۰۰۴)، بلینتون و آلن (۱۹۹۲) و السعید (۲۰۱۲) و همچنین در زمینه قابلیت اطمینان فازی افرادی چون کای و همکاران (۱۹۹۱)، چن (۱۹۹۴)، دکومن (۱۹۹۶) و هانگ (۱۹۹۵) تحقیقات گسترده‌ای انجام داده‌اند.

در این مقاله عملکرد سیستم زمانی که اطلاعات به صورت زبانی یا فازی هستند، ارزیابی می‌شود. مثالی در همین رابطه وجود دارد که می‌توان شبکه حسگر بی‌سیم را در نظر گرفت که شامل صدها یا هزاران حسگر با محدودیت‌های محاسباتی، انرژی و حافظه است که وظیفه‌ی دریافت اطلاعات از محیط پیرامون خود، آنالیز و پردازش داده‌ها و نیز ارسال داده‌های حس شده به دیگر گره‌ها و یا ایستگاه پایه را بر عهده دارد. در این نوع شبکه‌ها، گره‌های حسگر برای تامین انرژی خود به باتری‌هایی با توان اندک وابسته‌اند. از آنجایی که در این نوع شبکه‌ها، انرژی به‌عنوان یک مساله چالش برانگیز محسوب می‌گردد، پس باید قبل از خرابی باتری‌هایی که دارای طول عمر فازی هستند، سیستم هشدار فعال شود که باتری با توان بیشتر جایگزین این باتری شود.

۲ تابع قابلیت اطمینان

قابلیت اطمینان به احتمال فعال بودن یک سیستم منسجم گفته می‌شود که ارتباط مستقیمی با مؤلفه‌های آن دارد؛ قابلیت اطمینان یک سیستم به‌عنوان تابعی از قابلیت اطمینان اجزای آن مورد بحث قرار می‌گیرد. همچنین، قابلیت اطمینان را می‌توان به‌عنوان تابعی از زمان تعریف کرد. فرض کنید یک سیستم یا قطعه، دارای طول عمر T است. طبیعی است که T یک متغیر تصادفی است که می‌تواند در بازه $[0, \infty)$ مقدار بگیرد. متغیر T می‌تواند پیوسته و یا گسسته باشد، به عبارت دیگر سیستم می‌تواند طول عمری داشته باشد که هر مقدار مثبت را اختیار کند، یا می‌تواند طول عمر آن را طوری تعریف کرد که مقادیر شمارش‌پذیر را در بازه $[0, \infty)$ اختیار کند (۱). در این مقاله متغیرهای طول عمر پیوسته مد نظر است.

تعریف ۱.۲. (۱): فرض کنید متغیر تصادفی طول عمر T (به‌طور مطلق) پیوسته بوده و دارای تابع توزیع احتمال F و تابع چگالی f باشد، آنگاه تابع قابلیت اطمینان آن به‌صورت زیر تعریف می‌شود:

$$R(t) = P(T > t) = 1 - F(t) = \int_t^{\infty} f(x)dx. \quad (1)$$

۳ تابع دسترس‌پذیری

یک سؤال مهم درباره سیستمی که پس از شکست مورد تعمیر قرار گرفته آن است که، احتمال فعال بودن سیستم در زمان t چقدر است؟ در این رابطه کمیت مهمی به نام «دسترس‌پذیری» ارائه می‌شود که مبحث مهمی در تعمیر و نگهداری سیستم‌ها محسوب می‌شود. در واقع محاسبه مدت زمان عملیاتی بودن یک سیستم همان دسترس‌پذیری است. برای تعریف دسترس‌پذیری، متغیر $X(t)$ را به‌صورت زیر در نظر بگیرید:

$$X(t) = \begin{cases} 1 & , \text{ اگر سیستم در زمان } t \text{ کار کند} \\ 0 & , \text{ در غیر این صورت} \end{cases}$$

تعریف ۱.۳. (۴): دسترس‌پذیری آنی^۱ سیستم در زمان t ، که با $A(t)$ نمایش می‌دهیم، برابر است با:

$$A(t) = P(X(t) = 1) = E(X(t))$$

برای تعریف دسترس‌پذیری سیستم که یکی از شاخص‌های قابلیت اطمینان است، ابتدا بهتر است الگوی تعمیر سیستم‌ها را معرفی کنیم. در مباحث قابلیت اطمینان الگوهای احتمالی برای تعمیر و نگهداری سیستم‌ها وجود دارد. با توجه به اینکه در سیستم‌هایی که امروزه مورد استفاده قرار می‌گیرند، قطعات گران‌قیمت وجود دارد، پس به محض از کار افتادن یکی از قطعات سیستم توجیه اقتصادی برای عدم استفاده از آن سیستم وجود ندارد. بنابراین در چنین موقعیت‌هایی بایستی سیستم تعمیر و دوباره به وضعیت فعال اولیه برگردد. توجه شود که هر سیستمی این توانایی را ندارد و اگر سیستمی این قابلیت را داشته باشد به آن سیستم تعمیرپذیر می‌گوییم. برای چنین سیستم‌هایی تعیین الگوی احتمالی بهینه برای زمان تعویض و یا تعمیر قطعات سیستم‌ها حائز اهمیت می‌باشد.

^۱Instantaneous Availability

۱.۳ الگوی تعمیر کامل

در این الگو فرض کنید سیستم در زمان $t = 0$ شروع به کار می کند و در مدت زمان تصادفی X_1 فعال است و بعد از آن خراب می شود. همچنین تعمیر، Y_1 واحد زمانی به طول می انجامد. بعد از تعمیر اول سیستم مانند یک سیستم نو دوباره شروع به کار می کند و پس از X_2 واحد زمانی از کار می افتد، مدت زمان تعمیر نیز Y_2 واحد زمانی به طول می انجامد و روند به همین صورت ادامه دارد. در این روش اگر X_i ها متغیرهای تصادفی مستقل و هم توزیع با میانگین μ و Y_i ها متغیرهای تصادفی مستقل و هم توزیع با میانگین ν باشند، آنگاه $\{Z_i = X_i + Y_i, i \geq 1\}$ یک فرآیند تجدید است و متوسط زمان هر فرآیند تجدید برابر $\mu + \nu$ می باشد. در این صورت با این فرض که X_i و Y_i از هم مستقل و به ترتیب دارای توابع توزیع F و G باشند می توانیم تابع توزیع Z_i را به صورت زیر محاسبه کنیم:

$$\begin{aligned} F_Z(z) &= P(Z_i \leq z) = P(X_i + Y_i \leq z) \\ &= \int_0^z P(Y_i \leq z - x) f_{X_i}(x) dx \\ &= \int_0^z P(X_i \leq z - y) f_{Y_i}(y) dy. \end{aligned}$$

با توجه به مطالب فوق تعداد از کار افتادگی های سیستم در بازه $(0, t)$ را می توان به شکل زیر تعریف نمود:

$$N(0) = 0, \quad N(t) = \max \{n, S_n = \sum_{i=1}^n Z_i \leq t\}.$$

میانگین تعداد از کار افتادگی های سیستم را با $m(t)$ نشان می دهیم و به آن تابع چگالی تجدید^۲ گوئیم؛ یعنی،

$$m(t) = E(N(t)).$$

در ادامه محاسبه دسترس پذیری سیستم را با استفاده از تابع چگالی تجدید بیان می کنیم.

مفهوم دسترس پذیری که قبلاً در مورد آن صحبت شد، در روش تعمیر کامل با مدت زمان تصادفی تعمیر قابل تعریف است که برابر است با:

$$A(t) = R(t) + \int_0^t R(t-x) m'(x) dx, \quad (2)$$

که در آن $m'(x)$ مشتق تابع چگالی تجدید است (برای بهینه سازی زمان تجدید از مشتق تابع تجدید استفاده می شود). از رابطه (۲) دسترس پذیری آبی سیستم محاسبه می شود (۴). برای محاسبه دسترس پذیری علاوه بر تبدیلات مختلف از جمله تبدیل لاپلاس^۳، از روش هایی مانند زنجیرهای مارکوف نیز می شود استفاده کرد.

۴ محاسبات فازی

برای توسعه ارزیابی عملکرد سیستم ها و محاسبه دسترس پذیری سیستم های مهندسی فرض شده است که همه مشاهدات فازی (که در واقع همان طول عمر قطعات هستند)، به صورت اعداد فازی مثلثی در نظر گرفته شوند. این فرض علاوه بر افزایش دقت در ارزیابی دسترس پذیری، انجام محاسبات بر روی مجموعه های فازی را نیز ساده می سازد.

تعریف ۱.۴ (۳): عدد فازی $\tilde{A}$ که به صورت سه تایی مرتب $\tilde{A} = (a, b, c)$ نمایش داده می شود را فازی مثلثی گوئیم هرگاه، تابع درجه عضویت آن به صورت زیر مشخص شود:

$$\tilde{A}(x) = \begin{cases} \frac{x-a}{b-a}, & a \leq x < b \\ \frac{c-x}{c-b}, & b \leq x < c \\ 0, & \text{در غیر این صورت} \end{cases}$$

^۲Renewal Density Function

^۳Laplace Transform

تعریف ۲.۴. (۳): برش α عدد فازی مثلثی $\tilde{A}$ مجموعه مقادیری در $\tilde{A}$ را نشان می‌دهد که درجه عضویت آن‌ها بزرگتر یا مساوی α است و با $\tilde{A}[\alpha]$ نشان داده می‌شود؛ یعنی:

$$\tilde{A}[\alpha] = \{x \in X \mid \tilde{A}(x) \geq \alpha\} = [A^-(\alpha), A^+(\alpha)]$$

مراحل محاسبه‌ی دسترس‌پذیری فازی عبارت است از:

مرحله اول: انتخاب یک نمونه طول عمر از یک سیستم به صورت مشاهدات فازی

مرحله دوم: برآورد توزیع مشاهدات و پارامتر توزیع با استفاده از برش‌های α

مرحله سوم: محاسبه قابلیت اطمینان و تابع چگالی تجدید فازی برای برش‌های مختلف α

مرحله چهارم: محاسبه دسترس‌پذیری فازی و مقایسه آن با دسترس‌پذیری قطعی.

این مراحل در ادامه تشریح شده است:

مرحله اول: در این مرحله باید طول عمرهای فازی $\tilde{x}_i$ از یک توزیع خاص تولید شود؛ برای محاسبه دسترس‌پذیری سیستم تعداد ۱۰ داده فازی مثلثی از توزیع نمایی تولید شده که

به صورت زیر است:

$$\tilde{x}_1 = (1/763, 2/124, 2/485) \quad , \quad \tilde{x}_2 = (2/358, 2/445, 2/533)$$

$$\tilde{x}_3 = (2/211, 2/560, 2/908) \quad , \quad \tilde{x}_4 = (1/981, 2/608, 3/236)$$

$$\tilde{x}_5 = (3/408, 3/561, 3/714) \quad , \quad \tilde{x}_6 = (3/431, 4/733, 4/035)$$

$$\tilde{x}_7 = (3/500, 4/073, 4/646) \quad , \quad \tilde{x}_8 = (4/150, 4/713, 5/275)$$

$$\tilde{x}_9 = (4/671, 4/911, 5/151) \quad , \quad \tilde{x}_{10} = (4/964, 5/237, 5/510)$$

مرحله دوم: برای برآورد پارامتر توزیع نمایی با استفاده از برش‌های α از آن جایی که:

$$\tilde{x}_i[\alpha] = [x_i^-(\alpha), x_i^+(\alpha)]$$

و

$$\tilde{x}[\alpha] = \frac{1}{n} \sum \tilde{x}_i[\alpha] = \left[\frac{1}{n} \sum x_i^-(\alpha), \frac{1}{n} \sum x_i^+(\alpha) \right]$$

همچنین می‌دانیم:

$$\tilde{\lambda} = \frac{1}{\tilde{x}}$$

بنابراین:

$$\tilde{\bar{x}} = (3/243, 3/596, 3/949)$$

در نتیجه:

$$\tilde{\bar{x}}^{-1}(x) = \begin{cases} \frac{3/949x-1}{0/353x} & , \quad 3/949 \leq x < 3/596 \\ \frac{1-3/243x}{0/353x} & , \quad 3/596 \leq x < 3/243 \end{cases}$$

و

$$\tilde{\lambda} = (0/253, 0/278, 0/308)$$

مرحله سوم: از آن جایی که در این مقاله تعمیر از نوع کامل است بنابراین بعد از هر خرابی فرآیند تجدید رخ می‌دهد که تابع چگالی تجدید سیستم برای توزیع نمایی برابر است با:

$$\tilde{m}(t) = \{\lambda t \mid \lambda \in \tilde{\lambda}[\alpha]\}$$

$$\tilde{\lambda}[\alpha] = [\tilde{\lambda}^-(\alpha), \tilde{\lambda}^+(\alpha)] = \left[\frac{1}{3/949 - 0/353\alpha}, \frac{1}{3/243 + 0/353\alpha} \right]$$

قابلیت اطمینان فازی سیستم برابر است با:

$$\begin{aligned}\tilde{R}(t)[\alpha] &= \left\{ \int_t^\infty f(x)dx \mid \lambda \in \tilde{\lambda}[\alpha] \right\} \\ &= \{e^{-\lambda t} \mid \lambda \in \tilde{\lambda}[\alpha]\} \\ &= [e^{-\frac{1}{3.449-0.453\alpha}t}, e^{-\frac{1}{3.443+0.453\alpha}t}]\end{aligned}$$

که برای α برش‌های مختلف قابلیت اطمینان از جدول ۱ محاسبه می‌شود. حال قابلیت اطمینان قطعی سیستم برای داده‌های غیرفازی را محاسبه می‌کنیم:

جدول ۱: قابلیت اطمینان فازی سیستم

α	$\tilde{R}(t)[\alpha]$
۰	$[e^{-0.2532t}, e^{-0.3083t}]$
۰/۲	$[e^{-0.2578t}, e^{-0.3017t}]$
۰/۴	$[e^{-0.2626t}, e^{-0.2954t}]$
۰/۶	$[e^{-0.2675t}, e^{-0.2894t}]$
۰/۸	$[e^{-0.2727t}, e^{-0.2836t}]$
۱	$[e^{-0.2780t}, e^{-0.2780t}]$

$$x_1 = 2/124, x_2 = 2/445, x_3 = 2/560, x_4 = 2/608, x_5 = 3/561$$

$$x_6 = 3/733, x_7 = 4/073, x_8 = 4/713, x_9 = 4/911, x_{10} = 5/237$$

بنابراین $\lambda = 0.278$ و $R(t) = e^{-0.278t}$ است.

مرحله چهارم: محاسبه‌ی دسترس‌پذیری آنی سیستم

در این مرحله محاسبه دسترس‌پذیری فازی و قطعی محاسبه و با هم مقایسه می‌گردد. همانطور که در ادامه مشاهده می‌شود؛ دقت در محاسبه، در روش فازی بسیار بیشتر از روش قطعی است.

$$\tilde{A}(t)[\alpha] = \tilde{R}(t)[\alpha] + \int_t^\infty \tilde{R}(t-x)[\alpha]m'(t)[\alpha]dx$$

دسترس‌پذیری قطعی نیز به صورت $A(t) = e^{-0.287t} - 0.999e^{-0.287t} + 0.999$ خواهد بود. در ادامه نمودارهای مربوط به دسترس‌پذیری قطعی و دسترس‌پذیری فازی در

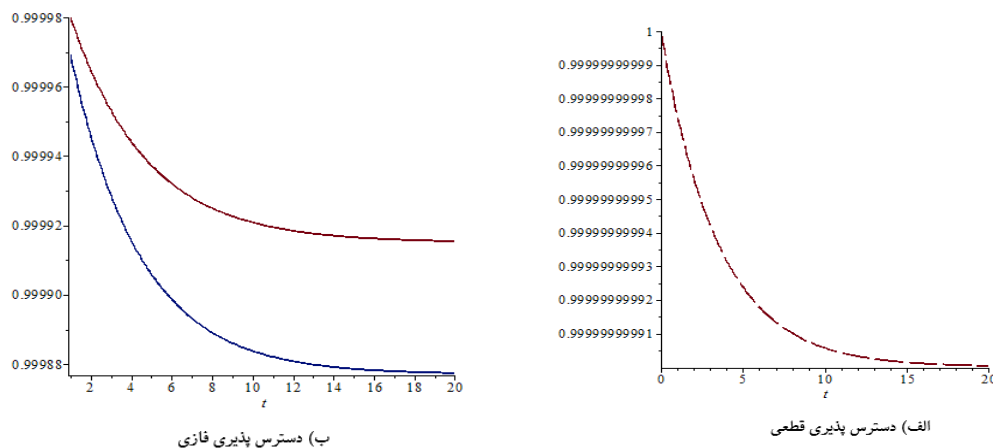
جدول ۲: دسترس‌پذیری فازی سیستم

α	$\tilde{A}(t)[\alpha]$
۰	$[0.0001132e^{-0.2532t} + 0.999, 0.0001831e^{-0.3083t} + 0.999]$
۰/۲	$[0.0002516e^{-0.2578t} + 0.999, 0.0001662e^{-0.3017t} + 0.999]$
۰/۴	$[0.0000192e^{-0.2626t} + 0.999, 0.0000071e^{-0.2954t} + 0.999]$
۰/۶	$[0.000085e^{-0.2675t} + 0.999, 0.000123e^{-0.2894t} + 0.999]$
۰/۸	$[0.000091e^{-0.2727t} + 0.999, 0.0000264e^{-0.2836t} + 0.999]$
۱	$[0.000034e^{-0.2780t} + 0.999, 0.000034e^{-0.2780t} + 0.999]$

شکل ۱ آورده شده است.

الف) $\alpha = 1$: در این حالت دسترس‌پذیری سیستم به ازای هر مقدار t خاص از روی نمودار قابل مشاهده است.

ب) $\alpha = 0.6$: در نمودار ب، دسترس‌پذیری با حداقل اطمینان 0.6 ، به ازای مقادیر مختلف t ، بین دو نمودار قرار می‌گیرد.



شکل ۱: نمودارهای مربوط به دسترس پذیری سیستم

همانطور که در نمودار الف و ب مشاهده می شود، دسترس پذیری سیستم در زمان $t = 8$ در حالت قطعی برابر است با 0.9991006607 و در حالت فازی با حداقل اطمینان 0.9999999999 و حداکثر اطمینان 0.9999999999 در بازه $[0.9999999999, 0.9999999999]$ قرار دارد.

۵ بحث و نتیجه گیری

با توجه به اهمیت موضوع قابلیت اطمینان در سیستم های مهندسی و پیش بینی شاخص های طول عمر و از طرفی وجود ابهام و داده های غیردقیق طول عمر موجب شده تا به منظور محاسبه شاخص های قابلیت اطمینان از مجموعه های فازی استفاده کنیم در این مقاله محاسبه قابلیت اطمینان و دسترس پذیری سیستم های مهندسی با طول عمر فازی که توزیع آنها نمایی می باشد، محاسبه شده است؛ همچنین می توان با محاسبه تابع چگالی تجدید برای هر توزیع دیگری مانند توزیع وایبل نیز دسترس پذیری را محاسبه نمود. از میان الگوهای مختلف تعمیر، در این مقاله الگوی تعمیر کامل در نظر گرفته شده است هر چند که برای تعمیر مینیمال و تعمیر ناقص نیز دسترس پذیری، با محاسبه تابع چگالی تجدید امکان پذیر است. برای تحقیقات آتی مرتبط با این موضوع نیز می توان سایر شاخص های طول عمر فازی و دسترس پذیری سیستم های k حالتی ($k > 2$) را محاسبه نمود.

مراجع

- [۱] اسدی، م. (۱۳۹۲). آشنایی با نظریه قابلیت اعتماد، مرکز نشر دانشگاهی.
- [۲] بلینتون، ر. و آلن، ر. ترجمه رضاییان، م. (۲۰۰۰). ارزیابی قابلیت اطمینان سیستم های مهندسی، مفاهیم و روش ها، انتشارات دانشگاه امیرکبیر، تهران.
- [۳] طاهری، م. و ماشین چی، م. (۱۳۸۷). مقدمه ای بر احتمال و آمار فازی، انتشارات دانشگاه شهید باهنر کرمان.
- [4] Barlow, R. E. and Proschan, F. (1975). *Statistical Theory of Reliability and Life Testing: probability models*. Florida State Univ Tallahassee.
- [5] Barlow, R. E. and Proschan, F. (1981). *Statistical Theory of Life Testing: Probability Models*. To Begin With, Silver Springs, MD.
- [6] Billinton, R. and Allan, R. N. (1992). *Reliability Evaluation of Engineering Systems*. New York: Plenum Press.
- [7] Chen, S. M. (1994). Fuzzy system reliability analysis using fuzzy number arithmetic operations. *Fuzzy sets and systems*, 64(1), 31-38.

- [8] Cui, L. and Li, j. (2004). Availability for a repairable system with finite repairs. *In Advanced Reliability Modeling* , 97-100.
- [9] De Cooman, G. (1996). On modeling possibilistic uncertainty in two-state reliability theory. *Fuzzy sets and systems*, 83(2), 215-238.
- [10] Elsayed, E. A. (2012). *Reliability Engineering* (Vol. 88). John Wiley & Sons.
- [11] Huang, H. Z. (1995). Reliability analysis method in the presence of fuzziness attached to operating time. *Microelectronics Reliability*, 35(12), 1483-1487.
- [12] Kai-Yuan, C., Chuan-Yuan, W., & Ming-Lian, Z. (1991). Fuzzy variables as a basis for a theory of fuzzy reliability in the possibility context. *Fuzzy sets and systems*, 42(2), 145-172.

تصمیم‌گیری متعادل در آزمون فرضیه‌های فازی

محبوبه سادات مدنی^۱، دانشگاه صنعتی شاهرود، دانشکده ریاضی، بخش آمار
محمد رضا ربیعی، دانشگاه صنعتی شاهرود، دانشکده ریاضی، بخش آمار
عباس پرچمی، دانشگاه شهید باهنر کرمان، دانشکده ریاضی و کامپیوتر، بخش آمار

چکیده

برای آزمون فرضیه مبتنی بر روش p -مقدار در حالت معمولی تنها فرضیه‌های H_0 در برابر H_1 را بررسی می‌کنیم. می‌دانیم p -مقدار مبتنی بر H_0 است و به H_1 بستگی ندارد. برای از بین بردن این نقطه ضعف روش p -مقدار جدید فازی متعادل را مطرح می‌کنیم. بعبارتی p -مقدار را علاوه بر فرضیه‌های H_0 در برابر H_1 ، برای فرضیه‌های H_1 در برابر H_0 نیز مورد بحث قرار می‌دهیم. هدف ما در این مقاله این است که ایده p -مقدار (۲) را در شیوه p -مقدار فازی (۵) از دو دیدگاه فرضیه‌های فازی H_0 در برابر H_1 و H_1 در برابر H_0 بررسی کنیم و شاخصی برای تصمیم‌گیری نهایی در محیط فازی مطرح کنیم. این روش p -مقدار فازی متعادل مزایایی نسبت به روش p -مقدار معمولی دارد که مهم‌ترین آن، تعادل در میزان بستگی تصمیم به هر دو فرضیه فازی H_0 و H_1 است.

کلمات کلیدی: آزمون فرضیه فازی، فرضیه‌های فازی، p -مقدار فازی متعادل.

۱ مقدمه

یک روش عملی برای آزمون فرضیه‌های H_0 در برابر H_1 ، مقایسه p -مقدار با سطح معنی‌داری است. در حالت معمولی، p -مقدار مبتنی بر H_0 محاسبه و H_1 نقش به‌سزایی در محاسبه آن ندارد. (۲) ایده p -مقدار پیشنهادی خود را، که از یک نگاه به H_0 و از نگاهی دیگر به H_1 توجه دارد معرفی کردند. در این روش p -مقدار علاوه بر آزمون فرضیه‌های H_0 در برابر H_1 ، مبتنی بر آزمون فرضیه‌های H_1 در برابر H_0 نیز است. تلاش‌های زیادی در زمینه آزمون فرضیه مبتنی بر p -مقدار توسط پژوهشگران آماری صورت گرفته است. از جمله (۳) آزمون فرضیه‌ها با داده‌های فازی را بررسی کردند. همچنین (۵؛ ۶) آزمون فرضیه‌های فازی با داده‌های دقیق و فازی را مطرح کردند. (۲) p -مقدار را علاوه بر فرضیه‌های H_0 در برابر H_1 ، برای فرضیه‌های H_1 در برابر H_0 مورد بحث قرار دادند. هدف ما در این مقاله این است که ایده p -مقدار عمادی و ارقامی را در شیوه p -مقدار فازی پرچمی و همکاران از دو دیدگاه H_0 در برابر H_1 و H_1 در برابر H_0 بررسی کنیم و شاخصی برای پذیرش (یا رد) فرضیات مطرح کنیم. در بخش دوم این مقاله آزمون فرضیه‌های آماری را بررسی می‌کنیم و در بخش سوم به انگیزه و تعریف فرضیه‌های فازی می‌پردازیم. بخش چهارم قاعده تصمیم‌گیری را بیان کرده و پس از بیان مثال برای روشن شدن ایده پیشنهادی، نتیجه‌گیری بیان می‌شود.

۲ آزمون فرضیه‌های آماری

فرض کنید $\mathbf{X} = (X_1, \dots, X_n)$ یک نمونه تصادفی با مقادیر مشاهده شده $\mathbf{x} = (x_1, \dots, x_n)$ باشد که X_i دارای تابع چگالی احتمال $f(x_i; \theta)$ ، $i = 1, \dots, n$ با پارامتر نامعلوم $\theta \in \Theta \subseteq \mathbb{R}$ باشد. مساله آزمون فرضیه‌های آماری تصمیم‌گیری برای پذیرش (یا رد) فرضیه $H_0: \theta \in \Theta_0 \subseteq \Theta$ در برابر $H_1: \theta \in \Theta_0^c = \Theta - \Theta_0$ بر اساس نمونه تصادفی $\mathbf{X}$ است. معمولاً فرضیه‌های آماری به یکی از صورت‌های جدول ۱ هستند که θ_1 و θ_0 دو عدد معلوم هستند و به ترتیب مرز فرضیه صفر و مرز فرضیه جایگزین نامیده می‌شوند. آزمون Φ آزمون در سطح $\alpha \in [0, 1]$ نامیده می‌شود اگر $\alpha_\Phi \leq \alpha$ که در آن $\alpha_\Phi = \sup_{\theta \in \Theta_0} P_\theta(RH_0)$. معمولاً آزمون آماری بر اساس آماره آزمون $T(X)$ است. در آزمون غیر تصادفی فضای مقادیر ممکن برای آماره آزمون T به ناحیه رد و مکمل آن ناحیه پذیرش تجزیه می‌شود. تحت شرایطی خاص ناحیه رد معمولاً به صورت زیر هستند

$$(۱) \quad \begin{array}{lll} \text{الف)} & T \leq t_l & \text{ب)} \quad T \geq t_r \\ \text{ج)} & T \notin (t_1, t_2) \end{array}$$

که t_1 و t_2 و t_l و t_r چندک‌های خاص توزیع T هستند.

۳ انگیزه و تعریف فرضیه‌های فازی

انگیزه معرفی فرضیه‌های فازی را به طور خلاصه با مثال زیر روشن می‌کنیم. لازم به ذکر است در این مقاله مجموعه‌های فازی را با حروف پر رنگ^۱ نمایش می‌دهیم.

مثال ۱.۳. فرض کنید یک آزمایشگر علاقمند به ارزیابی میانگین رشد گیاه با اندازه‌گیری قطر آن است که دارای توزیع نرمال با میانگین نامعلوم μ و انحراف استاندارد معلوم σ می‌باشد. روش معمول آزمون فرضیه $H_0: \mu = \mu_0$ در مقابل $H_1: \mu \neq \mu_0$ برای مقادیر معین μ_0 بر مبنای نمونه تصادفی $\mathbf{X}$ که $X_i \sim N(\mu, \sigma^2)$ ، $i = 1, \dots, n$ است، اما بدیهی است که اگر میانگین نمونه داده شده، کمی متفاوت‌تر از μ_0 باشد، آنگاه فرضیه صفر قابل قبول است. در صورتی که اندکی تفاوت μ_0 از μ ، سبب نقض معادله ریاضی مطرح شده در قالب فرضیه H_0 می‌شود. بنابراین مناسب‌تر و منطقی‌تر است فرضیه‌های H_0 و H_1 را به ترتیب با اصطلاحات فازی «نزدیک به μ_0 » و «دور از μ_0 » فرمول‌بندی کنیم. به عبارت دیگر فرضیه‌های واقع‌گرایانه‌تر «نزدیک به μ_0 است: H_0 » در مقابل «دور از μ_0 است: H_1 » آزمون شود.

در مواجه با چنین شرایط واقعی برخی فرضیه‌ها را به صورت «تابعی از H_0 است: H_0 » در مقابل «تابعی از H_1 است: H_1 »، صورت‌بندی می‌کنیم و آن را مساله آزمون فرضیه‌های فازی می‌نامیم. برای چنین مسائلی رویکرد p -مقدار فازی جدیدی را در این مقاله ارائه می‌دهیم که دارای مزایای مهمی از جمله هم‌زمانی بررسی هر دو فرضیه صفر و جایگزین است. ابتدا تعاریف فرضیه فازی و مرز فرضیه فازی که برگرفته از مراجع (۱) و (۵) است را بیان و از آن‌جا که نیاز به مقایسه p -مقدارهای فازی داریم ملاکی را برای مقایسه دو مقدار فازی بر اساس (۷) در تعریف ۲.۴ مطرح می‌کنیم و سپس به بیان تحلیل آزمون فرضیه‌های فازی به روش p -مقدار فازی متعادل می‌پردازیم.

تعریف ۲.۳. هر فرضیه به صورت «تابع H است: H » فرضیه فازی نامیده می‌شود که در آن «تابع H است: H » دلالت بر آن دارد که θ با درجه عضویت $H(\theta)$ متعلق به زیر مجموعه فازی H از فضای پارامتر Θ است.

تعریف ۳.۳. الف) فرضیه فازی «تابع H است: H » فرضیه فازی یکطرفه راست نامیده می‌شود اگر وجود داشته باشد $\theta_1 \in \Theta$ که $H(\theta) = 1$ برای $\theta \geq \theta_1$ ، و H یک تابع غیرنزولی از θ برای $\theta < \theta_1$ باشد. ب) فرضیه فازی «تابع H است: H » فرضیه فازی یکطرفه چپ نامیده می‌شود اگر وجود داشته باشد $\theta_1 \in \Theta$ که $H(\theta) = 1$ برای $\theta \leq \theta_1$ ، و H یک تابع غیرصعودی از θ برای $\theta > \theta_1$ باشد.

تعریف ۴.۳. مرز فرضیه فازی H یک زیرمجموعه فازی از Θ با تابع عضویت H_b بصورت زیر تعریف می‌شود

$$\begin{aligned} \text{الف)} \quad H_b(\theta) &= \begin{cases} H(\theta) & \theta \leq \theta_1 \\ 0 & \theta > \theta_1 \end{cases} \quad \text{اگر } H \text{ فرضیه فازی یکطرفه راست باشد.} \\ \text{ب)} \quad H_b(\theta) &= \begin{cases} H(\theta) & \theta \geq \theta_1 \\ 0 & \theta < \theta_1 \end{cases} \quad \text{اگر } H \text{ فرضیه فازی یکطرفه چپ باشد.} \end{aligned}$$

^۱ Bold

آزمون فرضیه‌های فازی

مشابه انواع فرضیه‌های معمولی بیان شده در ستون الف از جدول ۱، فرضیه‌های فازی نیز معمولاً به صورت ستون ب از جدول ۱ بیان می‌شوند که در آن θ_0 و θ_1 با استفاده از رابطه $(H_0)_\delta = [\theta_0(\delta), \theta_1(\delta)]$ تعریف می‌شوند. بدیهی است نواحی بحرانی فرضیه‌های فازی مشابه ناحیه بحرانی آزمون فرضیه‌های دقیق است. بعبارتی دیگر ناحیه بحرانی آزمون‌های ۱ و ۳ از جدول ۱ همانند بخش (الف) از رابطه (۱) است. همچنین ناحیه بحرانی آزمون‌های ۲ و ۴ از جدول ۱ همانند بخش (ب) از رابطه (۱) و ناحیه بحرانی آزمون ۵ از جدول ۱ همانند بخش (ج) از رابطه (۱) است.

جدول ۱: صورت کلی فرضیه‌های آماری H_0 و H_1 کلاسیک و فازی.

آزمون	فرضیه	الف) فرضیه‌های آماری کلاسیک	ب) فرضیه‌های آماری فازی
۱	H_0 H_1	$\theta = \theta_0$ $(\theta_0 > \theta_1) \theta = \theta_1$	θ تقریباً θ_0 است θ تقریباً θ_1 است $(\theta_0 > \theta_1)$
۲	H_0 H_1	$\theta = \theta_0$ $(\theta_0 < \theta_1) \theta = \theta_1$	θ تقریباً θ_0 است θ تقریباً θ_1 است $(\theta_0 < \theta_1)$
۳	H_0 H_1	$\theta \geq \theta_0$ $\theta < \theta_0$	θ تقریباً بزرگتر از θ_0 است θ تقریباً کوچکتر از θ_0 است
۴	H_0 H_1	$\theta \leq \theta_0$ $\theta > \theta_0$	θ تقریباً کوچکتر از θ_0 است θ تقریباً بزرگتر از θ_0 است
۵	H_0 H_1	$\theta = \theta_0$ $\theta \neq \theta_0$	θ نزدیک به θ_0 است θ دور از θ_0 است

۴ -p-مقدار فازی متعادل

قضیه ۰.۴. در مساله آزمون فرضیه‌های فازی H_0 در مقابل H_1 با داده‌های دقیق، (۵)، δ -برش‌های p -مقدار فازی فرضیه H_0 در مقابل H_1 ، در سه حالت مطرح شده در (۱) را به صورت زیر محاسبه کردند

$$P_\delta^{\circ 1} = [P_{\theta_1(\delta)}(T \leq t), P_{\theta_1(\delta)}(T \leq t)] \quad \text{الف)}$$

$$P_\delta^{\circ 1} = [P_{\theta_1(\delta)}(T \geq t), P_{\theta_1(\delta)}(T \geq t)] \quad \text{ب)}$$

$$P_\delta^{\circ 1} = \begin{cases} [P_{\theta_1(\delta)}(T \geq t), P_{\theta_1(\delta)}(T \geq t)] & \text{if } t \geq m_r \\ [P_{\theta_1(\delta)}(T \leq t), P_{\theta_1(\delta)}(T \leq t)] & \text{if } t \leq m_l \end{cases} \quad \text{ج)}$$

که در آن $\delta \in (0, 1)$ ، و توابع θ_0, θ_1 با استفاده از رابطه $(H_0)_\delta = [\theta_0(\delta), \theta_1(\delta)]$ حاصل می‌شوند. از آنجا که میانه آماره آزمون تحت مرز H_0 یک مجموعه فازی با تابع عضویت $m(m) = H_0(\theta)$ (که در آن میانه توزیع T تحت θ با m نشان داده شده) می‌باشد و لذا در حالت (ج) داریم $m_r = \sup\{m | m \in \text{supp}(m)\}$ ، $m_l = \inf\{m | m \in \text{supp}(m)\}$. با پیدا شدن δ -برش‌ها می‌توان p -مقدار فازی فرضیه H_0 در مقابل H_1 را بدست آورد که با نماد $P^{\circ 1}$ نمایش می‌دهیم.

در آمار کلاسیک، تصمیم در مورد فرضیه H_0 در مقابل H_1 ، بر اساس مقایسه p -مقدار مشاهده شده با α یعنی سطح معنی داری آزمون اتخاذ می‌گردد. اگر p -مقدار کمتر از α باشد، آن‌گاه فرضیه H_0 در سطح معنی داری α رد و درغیر اینصورت پذیرفته می‌شود. اشکال عمده این روش آن است که در محاسبه p -مقدار نقش فرضیه H_1 نادیده گرفته می‌شود و محاسبات در مورد آن صرفاً براساس H_0 صورت می‌گیرد. همچنین برای انجام آزمون‌های کلاسیک نیاز به تعیین سطح معنی داری α داریم و با توجه به عدم بودن فرمولی برای تعیین α ،

می‌توان ادعا کرد که تصمیم نهایی تابعی از α است که خود موجبات سوء استفاده از آزمون کلاسیک (برای موجه جلوه کردن نتایج) توسط آزمایشگر را نیز فراهم می‌سازد. به منظور رفع این نقایص، (۲) پیشنهاد به مقایسه $P^{\circ 1}$ با $P^{\circ 1}$ ، بجای مقایسه p -مقدار با α ، کردند. از آنجا که در این مقاله فرضیه‌ها و در نتیجه p -مقادیرها به صورت مجموعه‌های فازی هستند، باید این مقایسه به شکل فازی صورت گیرد. بنابراین نیاز به ملاکی برای مقایسه دو مجموعه فازی داریم. (۱) معیار D برای رتبه‌بندی دو مجموعه فازی که توسط (۷) ارائه شده بود به صورت تعریف ۲.۴ استفاده کردند.

تعریف ۲.۴. فرض کنید A و B دو مجموعه فازی (با توابع عضویت پیوسته) و

$$\Delta_{AB} = \int_{A_\delta^U > B_\delta^L} (A_\delta^U - B_\delta^L) d\delta + \int_{A_\delta^L > B_\delta^U} (A_\delta^L - B_\delta^U) d\delta \quad (2)$$

که در آن $A_\delta^U = \sup\{x | x \in A_\delta\}$ و $A_\delta^L = \inf\{x | x \in A_\delta\}$ ، به ترتیب انتها و ابتدای δ -برش مجموعه فازی A می‌باشند. به همین ترتیب B_δ^L و B_δ^U را می‌توان تعریف کرد. «درجه بزرگی A از B » به صورت زیر تعریف می‌گردد

$$D(A \succ B) = \frac{\Delta_{AB}}{\Delta_{AB} + \Delta_{BA}} \quad (3)$$

در این مقاله با تعمیم دیدگاه (۲)، قاعده تصمیم را به صورت زیر بیان می‌کنیم.

تعریف ۳.۴. در مساله آزمون فرضیه‌های فازی، $D(P^{\circ 1} \succ P^{\circ 1})$ را درجه پذیرش فرضیه فازی صفر نسبت به فرضیه فازی مقابل و $D(P^{\circ 1} \succ P^{\circ 1})$ را درجه پذیرش فرضیه فازی مقابل نسبت به فرضیه فازی صفر می‌نامیم، که در آن p -مقدار آزمون H در مقابل H_1 و p -مقدار آزمون H_1 در مقابل H است.

قاعده تصمیم‌گیری متعادل

در آزمون فرضیه‌های فازی، H_0 را در مقابل H_1 ، با درجه $D(P^{\circ 1} \succ P^{\circ 1})$ می‌پذیریم اگر $D(P^{\circ 1} \succ P^{\circ 1}) \geq D(P^{\circ 1} \succ P^{\circ 1})$ و در غیر این صورت H_1 را با درجه $D(P^{\circ 1} \succ P^{\circ 1})$ در مقابل H_0 می‌پذیریم.

مثال ۴.۴. طول عمر لامپ‌های تولیدی یک کارخانه (بر حسب ساعت) دارای توزیع نرمال با میانگین مجهول μ و انحراف معیار $\sigma = 120$ است. بر پایه مشاهدات نمونه‌ای تصادفی به حجم ۳۶ لامپ مقدار $\bar{x} = 1327$ ثبت شده است. قصد داریم فرضیه‌های فازی « H_1 : تقریباً 1400 است» در مقابل فرضیه فازی « H_0 : تقریباً 1300 است» را با توابع عضویت زیر آزمون کنیم

$$H_0(\mu) = \begin{cases} \frac{\mu - 1200}{100} & 1200 < \mu \leq 1300 \\ \frac{1400 - \mu}{100} & 1300 < \mu \leq 1400 \\ 0 & o.w \end{cases}$$

$$H_1(\mu) = \begin{cases} \frac{\mu - 1300}{100} & 1300 < \mu \leq 1400 \\ \frac{1500 - \mu}{100} & 1400 < \mu \leq 1500 \\ 0 & o.w \end{cases}$$

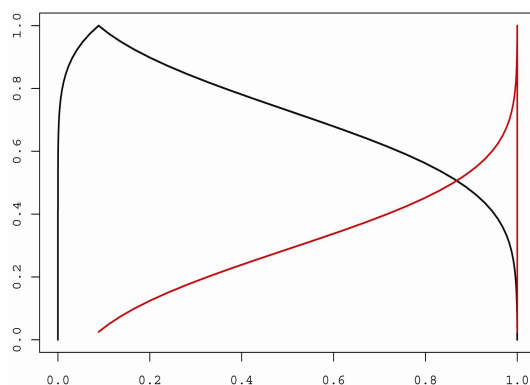
با استفاده از تعریف ۴.۳، $H_{jb}(\mu) = H_j(\mu)$ برای $j = 0, 1$. این مثال حالت خاصی از آزمون ۲ از جدول ۱ (ستون ب) است و ناحیه بحرانی برای آزمون H_0 در مقابل H_1 به صورت (ب) در رابطه (۱) است. با استفاده از قضیه ۱.۴، δ -برش‌های p -مقدار فازی H_0 در مقابل H_1 به صورت زیر به دست می‌آیند

$$P_\delta^{\circ 1} = \left[\int_{\frac{1200 - \mu_1(\delta)}{\sqrt{36}}}^{\infty} (\sqrt{\pi})^{-\frac{1}{2}} \exp\left(-\frac{z^2}{2}\right) dz, \int_{\frac{1200 - \mu_2(\delta)}{\sqrt{36}}}^{\infty} (\sqrt{\pi})^{-\frac{1}{2}} \exp\left(-\frac{z^2}{2}\right) dz \right]$$

که در آن به ازای هر $\delta \in (0, 1]$ داریم $\mu_1(\delta) = 100\delta + 1200$ ، $\mu_2(\delta) = 1400 - 100\delta$ با مشخص شدن کلیه δ -برش‌ها، p -مقدار فازی H_0 در مقابل H_1 مشخص می‌شود و نمودار تابع عضویت آن در شکل ۱ رسم شده است. از آنجا که علاوه بر $P^{\circ 1}$ به $P^{\circ 1}$ نیز نیاز داریم، که δ -برش‌های آن به صورت زیر محاسبه می‌شوند و نمودار تابع عضویت آن در

$$P_{\delta}^{\lambda^{\circ}} = \left[\int_{\frac{1327-\mu_2(\delta)}{\sqrt{36}}}^{\infty} (1 - (2\pi)^{-\frac{1}{2}} \exp(-\frac{z^2}{2})) dz, \int_{\frac{1327-\mu_1(\delta)}{\sqrt{36}}}^{\infty} (1 - (2\pi)^{-\frac{1}{2}} \exp(-\frac{z^2}{2})) dz \right]$$

پس از محاسبه $P_{\delta}^{\lambda^{\circ}}$ و $P_{\delta}^{\lambda^{\circ}}$ با استفاده از بسته نرم افزاری (۴)، $D(P^{\lambda^{\circ}} \succ P^{\lambda}) = 0.831$ و $D(P^{\lambda^{\circ}} \succ P^{\lambda^{\circ}}) = 0.869$ ، درجه پذیرش 0.831 پذیرفته می شود.



شکل ۱: توابع عضویت $P^{\lambda^{\circ}}$ (رنگ مشکی) و P^{λ} (رنگ قرمز) در مثال ۴.۲.

۵ نتیجه گیری

در این مقاله یک ایده مبتنی بر p -مقدار جدید فازی برای آزمون فرضیه های آماری زمانی که فرضیه ها فازی هستند ارائه شد. مزیت اصلی این ایده پیشنهادی مبتنی بر p -مقدار فازی، این است که بر اساس هر دو فرضیه صفر فازی و فرضیه جانشین فازی است. یک مثال برای تشریح روش پیشنهادی بیان شده است.

مراجع

- [۱] پرچی، ع. و ماشین چی، م. (۱۳۹۰). آزمون فرضیه در محیط فازی بر پایه p -مقدار، سری سیستم های فازی و محاسبات نرم، مباحثی در آمار و احتمال فازی، ۱۸-۱.
- [2] Emadi, M., and Arghami, N. (2003), Some measures of support for statistical hypotheses, J.Stat, Theory and Appl 2,165-176.
- [3] Filzmoser, P., and Viertl, R. (2004), Testing hypotheses with fuzzy data: the fuzzy p-value, Metrika 59, 21-29.
- [4] Parchami, A. (2017), FPV : Testing Hypotheses via Fuzzy P-value in Fuzzy Environment. R packege version 0.5.
- [5] Parchami, A., Taheri, S. M., and Mashinchi, M. (2010), Fuzzy p-value in testing fuzzy hypotheses with crisp data. *Statistical Papers* 51, 209-226.
- [6] Parchami, A., Taheri, S. M., Sadeghpour, B., and Mashinchi, M. (2016), A Simple but Efficient Approach for Testing Fuzzy Hypotheses. *Journal of Uncertainty Analysis and Applications*.
- [7] Yuan, Y. (1991), Criteria for evaluating fuzzy ranking methods, *Fuzzy Sets Syst* 43, 139-157.

غالب بودن شانس: معیاری برای مقایسه سیستم‌های پیچیده در فضای شانس

روح‌الله مهرعلی‌زاده^۱، محمد امینی^۲، بهرام صادقی‌پور گیلده^۲، حامد احمدزاده^۳

^{۱،۲} بخش آمار دانشگاه فردوسی مشهد

^۳ بخش علوم ریاضی دانشگاه سیستان و بلوچستان

چکیده

فضای شانس به عنوان فضای حاضضرب فضای احتمال و فضای نادقیق معرفی شده است. به عنوان دلیلی برای اهمیت معرفی این فضا می‌توان این‌طور بیان نمود که در بعضی مسائل ممکن است با سیستمی روبرو شویم که برخی از مولفه‌های آن متغیرهای تصادفی بوده در حالیکه برخی دیگر متغیرهای نادقیق می‌باشند. بنابراین برای مقایسه این چنین سیستمی در این مقاله مفهوم غالب بودن شانس معرفی شده و در ادامه نیز برخی از ویژگی‌های این معیار مورد بررسی قرار گرفته است. در نهایت کاربرد آن در انتخاب سبد سهام بهینه بطور ساده بیان شده است.

کلمات کلیدی: فضای شانس، متغیر تصادفی نادقیق، غالب بودن شانس، سیستم پیچیده و سبد سهام بهینه

۱ مقدمه

برای برآورد توزیع احتمال یک کمیت یا یک متغیر اگر تعداد نمونه کافی در اختیار داشته باشیم، نظریه احتمال روش‌های متنوع و موثری ارائه داده است. با این حال گاهی اوقات ممکن است نمونه در دسترس نباشد. در این مواقع باید از درجه باور متخصصان درباره کمیت مورد بررسی کمک گرفت. از آنجا که انسان به واسطه محافظه کار بودنش دامنه یک متغیر را معمولاً بیشتر از دامنه واقعی آن برآورد می‌کند ((۶))، بنابراین درجه باور را نمی‌توان به وسیله احتمال مدل نمود. به طوری که (۴) نشان داد که مدل کردن درجه باور بوسیله احتمال نه تنها مناسب نبوده، حتی ممکن است ما را به نتایج اشتباه برساند. از این‌رو برای برخورد با درجه باور (۱) نظریه نادقیقی را پیشنهاد نمود که این روزها، هم از جنبه تئوری و هم از جنبه کاربردی پیشرفت‌های قابل توجهی داشته است.

نظریه نادقیقی

تعریف ۱.۱. فرض کنید $\mathcal{L}$ یک σ -جبر روی مجموعه ناتهی Γ باشد. تابع مجموعه‌ای $[\circ, 1] \rightarrow \mathcal{M} : \mathcal{L}$ را یک اندازه نادقیق روی $\mathcal{L}$ گوئیم هرگاه در شرایط زیر صدق کند:

$$(1) \text{ به ازای هر مجموعه مرجع داشته باشیم } \mathcal{M}(\Gamma) = 1$$

$$(2) \text{ به ازای هر پیشامد دلخواه } \Lambda \text{ در } \mathcal{L} \text{ داشته باشیم } \mathcal{M}(\Lambda) + \mathcal{M}(\Lambda^c) = 1$$

$$(3) \text{ به ازای هر دنباله شمارش‌پذیر از پیشامدها داشته باشیم } \mathcal{M}(\bigcup_{i=1}^{\infty} \Lambda_i) \leq \sum_{i=1}^{\infty} \mathcal{M}(\Lambda_i)$$

$$(4) \text{ اگر به ازای هر } k = 1, 2, \dots \text{ سه تایی } (\Gamma_k, \mathcal{L}_k, \mathcal{M}_k) \text{ یک فضای نادقیق باشد، آنگاه به ازای هر پیشامد دلخواه } \Lambda_k \text{ در } \mathcal{L}_k \text{ داشته باشیم } \mathcal{M}(\prod_{i=1}^{\infty} \Lambda_i) = \bigwedge_{i=1}^{\infty} \mathcal{M}(\Lambda_i)$$

^۱ نام ارائه دهنده مقاله: ro.mehralizade@mail.um.ac.ir

تعریف ۲.۱. یک متغیر نادقیق τ تابعی است از فضای نادقیق $(\Gamma, \mathcal{L}, \mathcal{M})$ به زیرمجموعه‌ای از اعداد حقیقی به طوری که $\{\tau \in B\}$ به ازای هر مجموعه بورل B از اعداد حقیقی یک پیشامد خواهد بود.

تعریف ۳.۱. برای یک متغیر نادقیق τ تابع توزیع نادقیق آن را با نماد $\Upsilon_\tau(\cdot)$ نمایش داده و بصورت $\Upsilon_\tau(x) = \mathcal{M}(\tau \leq x)$ تعریف می‌شود.

تعریف ۴.۱. فرض کنید τ_1 و τ_2 به ترتیب دو متغیر نادقیق با توابع توزیع $\Upsilon_{\tau_1}(\cdot)$ و $\Upsilon_{\tau_2}(\cdot)$ باشند. اگر به ازای هر $t \in \mathbb{R}$ داشته باشیم

$$\Upsilon_{\tau_1}^{(k)}(t) \leq \Upsilon_{\tau_2}^{(k)}(t)$$

آنگاه گوییم متغیر نادقیق τ_1 متغیر نادقیق τ_2 را به طور نادقیق در مرتبه k ام مغلوب می‌کند و آن را با نماد $\tau_1 \succeq_{(k)} \tau_2$ نمایش می‌دهیم.

به خواننده برای آشنایی بیشتر با نظریه نادقیقی و همچنین ضرورت‌های معرفی آن منبع (۶) و برای آشنایی با کاربردهای آن منابع (۲) و (۳) پیشنهاد می‌شود. در عمل ممکن است با موضوعاتی روبرو شویم که بعضی از کمیت‌های آن‌ها با متغیرهای تصادفی و برخی دیگر از کمیت‌های آن‌ها با متغیرهای نادقیق مدل شوند. در این شرایط با سیستمی روبرو خواهیم بود که در آن تصادفی بودن و نادقیق بودن با هم اتفاق می‌افتد. به این سیستم در اصطلاح یک سیستم پیچیده گفته می‌شود. برای برخورد با یک سیستم پیچیده (۵) نظریه شانس را که به صورت حاصلضرب فضای احتمال و فضای نادقیق می‌باشد، را پیشنهاد نمود. نظریه شانس

تعریف ۵.۱. فرض کنید $(\Gamma, \mathcal{L}, \mathcal{M}) \times (\Omega, \mathcal{F}, \text{Pr})$ یک فضای شانس و $\Theta \in \mathcal{L} \times \mathcal{F}$ یک پیشامد تصادفی نادقیق باشد. آنگاه اندازه شانس پیشامد Θ به صورت زیر تعریف می‌شود.

$$\text{Ch}(\Theta) = \int_0^1 \text{Pr}\{\omega \in \Omega | \mathcal{M}\{\gamma \in \Gamma | (\omega, \gamma) \in \Theta\} \geq r\} dr$$

که دارای ویژگی‌های زیر می‌باشد

$$\text{Ch}(\Gamma \times \Omega) = 1 \quad (1)$$

$$\text{Ch}(\Theta) + \text{Ch}(\Theta^c) = 1 \quad \text{در } \mathcal{L} \times \mathcal{F} \text{ داریم} \quad (2)$$

$$\text{Ch}(\Theta_1) \leq \text{Ch}(\Theta_2) \quad \text{اگر داشته باشیم } \Theta_1 \subset \Theta_2 \quad \text{آنگاه داریم} \quad (3)$$

تعریف ۶.۱. یک متغیر تصادفی نادقیق ξ تابعی است اندازه‌پذیر از فضای شانس $(\Gamma, \mathcal{L}, \mathcal{M}) \times (\Omega, \mathcal{F}, \text{Pr})$ به زیرمجموعه‌ای از اعداد حقیقی به طوری که $\{\xi \in B\}$ به ازای هر مجموعه بورل B از اعداد حقیقی یک پیشامد خواهد بود.

تعریف ۷.۱. برای یک متغیر تصادفی نادقیق ξ تابع توزیع شانس آن را با نماد $\Phi_\xi(\cdot)$ نمایش داده و بصورت $\Phi_\xi(x) = \text{Ch}(\xi \leq x)$ تعریف می‌شود.

به خواننده برای آشنایی بیشتر با نظریه شانس منابع (۵) و (۶) پیشنهاد می‌شود.

در مباحث کاربردی چون اقتصاد، قابلیت اعتماد و بقا روبرو شدن با مسائلی که در آن‌ها تصادفی بودن و نادقیق بودن با هم وجود دارد، اجتناب ناپذیر خواهد بود. بنابراین برای مقایسه این سیستم‌های پیچیده نیاز داشتن به معیاری برای مقایسه ضروری به نظر می‌رسد. برای این منظور در این مقاله مفهوم غالب بودن شانس معرفی و برخی از ویژگی‌های آن بیان می‌شود. در ادامه نیز به کاربردی از این مفهوم اشاره خواهد شد.

۲ غالب بودن شانس: تعریف و ویژگی‌های آن

فرض کنید ξ یک متغیر تصادفی نادقیق با تابع توزیع $\Phi_\xi(\cdot)$ باشد. برای $k = 1, 2, \dots$ تعریف می‌کنیم

$$\Phi_\xi^{(k+1)}(t) = \int_{-\infty}^t \Phi_\xi^{(k)}(x) dx,$$

که در آن $\Phi_\xi^{(1)}(t) = \Phi_\xi(t)$ می‌باشد.

توجه ۱۰۲. در این مقاله متغیرهای تصادفی نادقیقی در نظر گرفته شده اند که شرایط زیر برای آنها برقرار می باشد.

$$\lim_{x \rightarrow -\infty} \Phi_{\xi}(x) = 0 \quad (۱) \quad \lim_{x \rightarrow +\infty} \Phi_{\xi}(x) = 1 \quad (۲) \quad \int_{-\infty}^{+\infty} x d\Phi_{\xi}(x) < +\infty \quad (۳)$$

تعریف ۲۰۲. فرض کنید ξ_1 و ξ_2 به ترتیب دو متغیر تصادفی نادقیقی با توابع توزیع $\Phi_{\xi_1}(\cdot)$ و $\Phi_{\xi_2}(\cdot)$ باشند. اگر به ازای هر $t \in \mathbb{R}$ داشته باشیم

$$\Phi_{\xi_1}^{(k)}(t) \leq \Phi_{\xi_2}^{(k)}(t)$$

آنگاه گوییم متغیر تصادفی نادقیقی ξ_1 متغیر تصادفی نادقیقی ξ_2 را به طور شانس در مرتبه k مغلوب می کند و آن را با نماد $\xi_1 \succeq_{(k)} \xi_2$ نمایش می دهیم.

توجه ۳۰۲. اگر داشته باشیم $\xi_2 \succeq_{(k)} \xi_1$ در حالیکه $\xi_1 \not\succeq_{(k)} \xi_2$ ، آنگاه گوییم متغیر تصادفی نادقیقی ξ_1 متغیر تصادفی نادقیقی ξ_2 را به طور شانس اکید در مرتبه k مغلوب می کند و آن را با نماد $\xi_2 \succeq_{(k)} \xi_1$ نمایش می دهیم.

توجه ۴۰۲. اگر در تعریف ۱۰۲ همه متغیرها متغیرهای تصادفی باشند، آنگاه این تعریف با تعریف غالب بودن تصادفی یکی خواهد شد و اگر تمام متغیرها متغیرهای نادقیقی باشند، آنگاه با تعریف غالب بودن نادقیقی یکی خواهد شد. بنابراین تعریف غالب بودن شانس را می توان به عنوان تعمیمی برای تعریف غالب بودن تصادفی و غالب بودن نادقیقی در نظر گرفت.

اگر برای $k = 1, 2, \dots$ فضای C_k را به صورت $C_k = \{\xi \mid E(|\xi|^k) \leq \infty\}$ تعریف کنیم که در آن $\|\xi\|_k = [E(|\xi|^k)]^{\frac{1}{k}}$ نرم این فضا می باشد، آنگاه با توجه به تعریف غالب بودن شانس (تعریف ۱۰۲) ویژگی های زیر به راحتی نتیجه می شود.

گزاره ۵۰۲. اگر متغیرهای تصادفی نادقیقی ξ_1 ، ξ_2 و ξ_3 در C_k باشند، آنگاه

$$(۱) \text{ به ازای هر } i = 1, 2, \dots, k, \xi_1 \succeq_{(i)} \xi_2.$$

$$(۲) \text{ به ازای هر } k, i = 1, 2, \dots, \text{ اگر داشته باشیم } \xi_2 \succeq_{(i)} \xi_1 \text{ و } \xi_3 \succeq_{(i)} \xi_2 \text{ آنگاه خواهیم داشت } \xi_3 \succeq_{(i)} \xi_1.$$

$$(۳) \text{ به ازای هر } k-1, i = 1, 2, \dots, \text{ اگر رابطه } \xi_2 \succeq_{(i)} \xi_1 \text{ برقرار باشد، آنگاه رابطه } \xi_2 \succeq_{(i+1)} \xi_1 \text{ نیز برقرار خواهد بود.}$$

در یک مقاله (۷) ترتیب شانس را به صورت زیر تعریف کرده و به عنوان کاربردی از آن در قابلیت اعتماد، به مقایسه سیستم های پیچیده k از n با استفاده از این ترتیب پرداخته اند.

تعریف ۶۰۲. اگر ξ_1 و ξ_2 دو متغیر تصادفی نادقیقی باشند، آنگاه گوییم ξ_2 در ترتیب شانس از ξ_1 کوچکتر است اگر به ازای هر $t \in \mathbb{R}$ داشته باشیم

$$\text{Ch}(\xi_2 > t) \leq \text{Ch}(\xi_1 > t)$$

و با نماد $\xi_2 \preceq_{ch} \xi_1$ نمایش داده می شود.

به عنوان ویژگی دیگر از تعریف غالب بودن شانس داریم

$$\text{گزاره ۷۰۲. اگر } \xi_1 \text{ و } \xi_2 \text{ دو متغیر تصادفی نادقیقی باشند، آنگاه } \xi_2 \succeq_{(1)} \xi_1 \text{ اگر و فقط اگر } \xi_2 \preceq_{ch} \xi_1.$$

اثبات.

$$\text{Ch}(\xi_1 > t) \leq \text{Ch}(\xi_2 > t) \iff \text{Ch}(\xi_1 \leq t) \geq \text{Ch}(\xi_2 \leq t) \iff \Phi_{\xi_1}(t) \geq \Phi_{\xi_2}(t)$$

□

۳ کاربرد

یکی از مباحث مهم مالی بحث سرمایه گذاری و خرید و فروش سهام می باشد. در بحث سرمایه گذاری، اگر امکانش فراهم باشد همواره انتخاب یک سبد سهام^۱ (مجموعه ای از چند محصول) به جای یک محصول عملکرد بهتری خواهد داشت. بنابراین زمانیکه با چند سبد سهام روبرو هستیم اینکه کدام سبد سهام را انتخاب کنیم تا به سود بالاتری برسیم، مهم خواهد بود.

^۲ برهان در صورت درخواست قابل ارائه می باشد.

^۱ Portfolio

برای یک سبد سهام اگر نرخ بازده برخی از اقلام آن متغیر باشد، آن را یک سبد سهام ریسکی گوئیم. اگر در یک سبد سهام ریسکی با n قلم، نرخ بازده برخی از اقلام متغیر نادقیق $(\tau_1, \tau_2, \dots, \tau_m)$ و نرخ بازده برخی دیگر متغیر تصادفی $(\eta_{m+1}, \eta_{m+2}, \dots, \eta_n)$ باشد، آنگاه میزان ثروت آتی (بازده شانسی) ما بعد از سرمایه‌گذاری یک متغیر تصادفی نادقیق (ξ) خواهد بود. با فرض اینکه بازده شانسی ما تابعی خطی بر حسب نرخ بازده‌های اقلام سبد سهام باشد می‌توان نوشت

$$\xi_w = w_1 \tau_1 + w_2 \tau_2 + \dots + w_m \tau_m + w_{m+1} \eta_{m+1} + w_{m+2} \eta_{m+2} + \dots + w_n \eta_n$$

که در آن w_i کسری از مقدار سرمایه‌گذاری شده برای سبد سهام مورد نظر می‌باشد که به قلم i ام تعلق گرفته است. هدف ماکزیم کردن بازده شانسی مورد انتظار یعنی $\max E(\xi_w)$ می‌باشد. اگر یک مقدار بازده شانسی ξ° در دسترس باشد، آنگاه مسئله انتخاب سبد سهام بهینه از حل دستگاه زیر حاصل می‌شود

$$\begin{cases} \max E(\xi_w) \\ \text{subject to:} \\ \xi_w \succeq_{(k)} \xi^\circ \\ w_1 + w_2 + \dots + w_n = 1 \\ w_i \geq 0 \quad i = 1, 2, \dots, n \end{cases}$$

برای حل این دستگاه می‌توان از روش برنامه ریزی خطی و در صورت پیچیده شدن مسئله از روشهای عددی کمک گرفت.

بحث و نتیجه‌گیری

در این مقاله مفهوم غالب بودن شانسی به عنوان تعمیمی از غالب بودن تصادفی و غالب بودن نادقیق بیان و سپس بعضی از ویژگی‌های غالب بودن شانسی ارائه شد. در نهایت به عنوان یک کاربرد، الگوی بهینه‌سازی ساده‌ای را برای انتخاب سبد سهام بهینه در مبحث سرمایه‌گذاری و انتخاب سبد سهام تشکیل دادیم که روش حل این الگو بسته به شرایط با استفاده از روشهای پیشرفته ریاضی نظیر برنامه‌ریزی خطی یا روشهای عددی میسر خواهد بود.

مراجع

- [1] Liu B. (2007), *Uncertainty theory*, 2nd ed., Berlin, Springer-Verlag.
- [2] Liu B. (2009), *Theory and practice of uncertain programming*, 2nd ed., Berlin, Springer-Verlag.
- [3] Liu B. (2010), *Uncertainty theory: A branch of mathematics for modeling human uncertainty*, Berlin, Springer-Verlag.
- [4] Liu B. (2012), Why is there a need for uncertainty theory?, *Journal of Uncertain Systems*, 6, 1, 3-10.
- [5] Liu Y. (2013), Uncertain random variables: a mixture of uncertainty and randomness, *Soft Computing*, 17, 4, 625-634.
- [6] Liu B. (2015), *Uncertainty theory*, 4th ed., Berlin, Springer-Verlag.
- [7] Mehralizade R., Amini M., Sadeghpour Gildeh B., and Ahmadzade H. (2018), Chance order of two uncertain random variables, *Journal of Uncertain Systems*, In Press.

Fuzzy process capability indices for simple linear profile

Zainab Abbasi Ganji¹ , Bahram Sadeghpour Gildeh²,

¹Khorasan Razavi Agricultural and Natural Resources Research and Education Center, AREEO, Mashhad, Iran

²Department of Statistics, Faculty of Mathematical Sciences, Ferdowsi University of Mashhad, Mashhad, Iran

Abstract

Process capability indices are numerical tools that quantify how well a process can meet customer requirements, specifications or engineering tolerances. In real world, in some cases we are faced with imprecise data. So, fuzzy logic is engaged to deal with them. This paper develops two fuzzy methods for measuring the process capability in simple linear profiles in the circumstances in which lower and upper specification limits are imprecise. Furthermore, numerical example is presented to illustrate the applicability of the proposed methods.

Keywords: Simple linear profile, Process capability index, Fuzzy logic, Triangular fuzzy numbers, Ranking function.

1 Introduction

Process capability indices (PCIs) are numerical quantities that measure how the process products meet the required specifications considered for them. In some processes, the quality of a process or product can be characterized by a functional relationship between a response quality characteristic and one or more explanatory variables. This relationship is usually known as a profile. Quality profiles can take on several different functional structures such as linear, nonlinear, polynomial, etc. Most research on profile monitoring are related to simple linear profile. A simple linear profile is a linear profile with only one independent variable.

(2) proposed two methods for measuring the process capability in simple linear profiles. The first method uses the percentage of nonconforming parts produced at each level of the independent variable to introduce a PCI. In the second method, based on the estimated responses, a vector of three components is defined to assess the process capability. In the next section, we refer to these indices.

¹Speaker: z.ganji@areeo.ac.ir, abbasiganji@mail.um.ac.ir

In real applications, there are processes that their specification limits are imprecise or in linguistic form, so traditional approaches may fail to account the capability of these processes. A graded representation of the specification limits is handled by employing the theory of fuzzy sets. In addition, fuzzy set theory is used to add more information and flexibility to PCIs than the traditional ones, but there is no research dedicated to the profile capability analysis in fuzzy environment. This paper aims to provide a research of the application of fuzzy logic in profile capability indices.

2 Process capability indices for simple linear profile

Suppose the relationship between the response variable and explanatory variable is such a simple linear regression model. Let (x_i, y_{ij}) ; $i = 1, 2, \dots, k$ be the observations for the j^{th} sample collected over time. When the process is in statistical control, the relationship between the response, independent variable and error terms can be modeled by $Y_{ij} = A_0 + A_1 X_i + \varepsilon_{ij}$; $i = 1, 2, \dots, k$, $j = 1, 2, \dots, n$, where ε_{ij} 's are independently and identically distributed normal random variables with mean zero and variance σ^2 . We can estimate A_0 and A_1 by a_0 and a_1 according to the sample observations as $\hat{A}_0 = a_0 = \frac{\sum_{j=1}^n a_{0j}}{n}$ and $\hat{A}_1 = a_1 = \frac{\sum_{j=1}^n a_{1j}}{n}$, where $a_{0j} = \bar{y}_j - a_{1j}\bar{x}$ and $a_{1j} = \frac{S_{xy(j)}}{S_{xx}}$, and $\bar{y}_j = \sum_{i=1}^k y_{ij}/k$, $\bar{x} = \sum_{i=1}^k x_i/k$, $S_{xy(j)} = \sum_{i=1}^k (x_i - \bar{x})y_{ij}$, and $S_{xx} = \sum_{i=1}^k (x_i - \bar{x})^2$; see (7).

Hence, we can write $\hat{Y}_{ij} = a_{0j} + a_{1j}x_i$, $i = 1, 2, \dots, k$, where $\hat{Y}_{ij}$ denotes the predicted value of response variable from j^{th} sample for a given level of X . Furthermore, $\hat{Y}_i = a_0 + a_1x_i$ is the predicted value of response variable for a given level of the independent variable.

The unbiased estimator of σ^2 can be defined by $MSE = \sum_{j=1}^n \sum_{i=1}^k e_{ij}^2 / n(n-2)$, and $e_{ij} = y_{ij} - \hat{y}_{ij}$. e_{ij} s are the residuals for the j^{th} sample. For more information, we refer to (1).

2.1 Capability assessment using responses

Let LSL_i and USL_i are the lower and upper specification limits of the response variable at the i^{th} level of independent variable. The average of the percentage of non-conforming profile produced can be defined by $P_L = \frac{\sum_{i=1}^k p_{Li}}{k}$ and $P_U = \frac{\sum_{i=1}^k p_{Ui}}{k}$, where, $p_{Li} = P_r(Y_i < LSL_i) = \Phi\left(\frac{LSL_i - \hat{Y}_i}{\hat{\sigma}_{Y_i}}\right)$ and $p_{Ui} = P_r(Y_i > USL_i) = 1 - \Phi\left(\frac{USL_i - \hat{Y}_i}{\hat{\sigma}_{Y_i}}\right)$, and $\Phi(\vartheta)$ is the ϑ percentile of standard normal distribution.

Process capability index is defined as;

$$\hat{C}_{pkA} = \frac{\min\{Z_L, Z_U\}}{3}, \quad (1)$$

where, $Z_L = Z(P_L) = \Phi^{-1}(1 - P_L)$ and $Z_U = Z(P_U) = \Phi^{-1}(1 - P_U)$.

2.2 Capability assessment using predicted responses

As it mentioned by (4), the linear fit or the prediction $\hat{Y}_i$ is distributed normally with mean $A_0 + A_1x_i$ and variance $\sigma^2\left(\frac{1}{k} + \frac{(x_i - \bar{x})^2}{S_{xx}}\right)$. In addition, $\hat{\mathbf{Y}} = (\hat{Y}_1, \hat{Y}_2, \dots, \hat{Y}_k)'$ has the multivariate normal distribution with the variance-covariance matrix as $\Sigma_{\hat{\mathbf{Y}}} =$

$\mathbf{W}(\mathbf{W}'\mathbf{W})^{-1}\mathbf{W}'\sigma^2$, where

$$\mathbf{W} = \begin{pmatrix} 1 & x_1 \\ 1 & x_2 \\ \vdots & \\ 1 & x_k \end{pmatrix}.$$

(2) proposed multivariate approaches of process capability to determine the process capability index for a profile based on $\widehat{Y}_i$ s and selected the method proposed by (9). The process capability vector has three components as $[C_{pM}, PV, LI]$. The first component, C_{pM} is defined by

$$C_{pM} = \left[\frac{\prod_{i=1}^k (USL_i - LSL_i)}{\prod_{i=1}^k (UPL_i - LPL_i)} \right]^{\frac{1}{k}}, \quad (2)$$

where, $LPL_i = \mu_i - \sqrt{\chi_{k,0.0027}^2} \sigma_i$ and $UPL_i = \mu_i + \sqrt{\chi_{k,0.0027}^2} \sigma_i$, and $\chi_{k,p}^2$ is the upper p quantile of a chi-square distribution with k degrees of freedom and μ_i and σ_i are the mean and standard deviation of $\widehat{Y}_i$, respectively. For more details, see (5). When the value of C_{pM} is greater than 1, it indicates that the circumscribed “modified process region” is smaller than the engineering specified region; namely, the part is “acceptable”.

The second component of this vector is the significant level of the observed value with the Hotelling's T^2 statistic, $t^2 = n(\widehat{\mathbf{Y}} - \boldsymbol{\mu}_0)' \boldsymbol{\Sigma}_{\widehat{\mathbf{Y}}}^{-1} (\widehat{\mathbf{Y}} - \boldsymbol{\mu}_0)$, for testing the process centrality, expressed as $PV = P_r(F_{k,n-k} > \frac{n-k}{k(n-1)} t^2)$, where $F_{k,n-k}$ denotes a random variable having the F distribution with $k, n-k$ degrees of freedom and $\boldsymbol{\mu}_0 = (\mu_{01}, \mu_{02}, \dots, \mu_{0k})'$ in which $\mu_{0i} = (LSL_i + USL_i)/2$. The third component (LI), taking values, 0 or 1, compares the location of the modified process region and the tolerance region. It takes the value 1 when the entire modified process region is contained within the tolerance region; otherwise, it takes the value 0. For more information, we refer to (8).

3 Fuzzy capability indices for simple linear profile

Here, we consider a situation in which the specification limits and target are triangular fuzzy numbers as $\widetilde{LSL} = T(l_1, l_2, l_3)$ and $\widetilde{USL} = T(u_1, u_2, u_3)$, respectively. It should be noted that in this paper, to make a comparison between two fuzzy numbers (quantities), we employ the ranking method proposed by (3).

3.1 Fuzzy estimation of the index $\widehat{C}_{pkA}$

First we estimate P_L , P_U , Z_L and Z_U according to the fuzzy set theory. Fuzzy estimations of P_L and P_U can be obtained as;

$$\widetilde{P}_L = \frac{\bigoplus_{i=1}^k \widetilde{p}_{L_i}}{k}, \quad \widetilde{P}_U = \frac{\bigoplus_{i=1}^k \widetilde{p}_{U_i}}{k}. \quad (3)$$

The indices $\widetilde{p}_{L_i}$ and $\widetilde{p}_{U_i}$ are gain by palcing their α -cut intervals one on top of the other, which have the form:

$$\widetilde{p}_{L_i}(\alpha) = [\widetilde{p}_{L_{il}}(\alpha), \widetilde{p}_{L_{ir}}(\alpha)], \quad \widetilde{p}_{U_i}(\alpha) = [\widetilde{p}_{U_{il}}(\alpha), \widetilde{p}_{U_{ir}}(\alpha)], \quad (4)$$

where $\widetilde{p}_{L_{il}}(\alpha) = P_r(Y_i < \widetilde{LSL}_{il}(\alpha))$, $\widetilde{p}_{L_{ir}}(\alpha) = P_r(Y_i < \widetilde{LSL}_{ir}(\alpha))$, $\widetilde{p}_{U_{il}}(\alpha) = P_r(Y_i > \widetilde{USL}_{il}(\alpha))$ and $\widetilde{p}_{U_{ir}}(\alpha) = P_r(Y_i > \widetilde{USL}_{ir}(\alpha))$. Therefore, α -cut intervals of $\widetilde{P}_L$ and $\widetilde{P}_U$ are obtained from the α -cut intervals of $\widetilde{p}_{L_i}$ and $\widetilde{p}_{U_i}$ according to the following

equations;

$$\begin{aligned}\tilde{P}_L(\alpha) &= [\tilde{P}_{L_l}(\alpha), \tilde{P}_{L_r}(\alpha)] = \left[\frac{\sum_{i=1}^k \tilde{p}_{L_{il}}(\alpha)}{k}, \frac{\sum_{i=1}^k \tilde{p}_{L_{ir}}(\alpha)}{k} \right], \\ \tilde{P}_U(\alpha) &= [\tilde{P}_{U_l}(\alpha), \tilde{P}_{U_r}(\alpha)] = \left[\frac{\sum_{i=1}^k \tilde{p}_{U_{il}}(\alpha)}{k}, \frac{\sum_{i=1}^k \tilde{p}_{U_{ir}}(\alpha)}{k} \right].\end{aligned}\quad (5)$$

Let $\tilde{Z}_L = \Phi^{-1}(1 \ominus \tilde{P}_L)$ and $\tilde{Z}_U = \Phi^{-1}(1 \ominus \tilde{P}_U)$, so these fuzzy quantities α -cut intervals can be obtained by

$$\begin{aligned}\tilde{Z}_L(\alpha) &= [\tilde{Z}_{L_l}(\alpha), \tilde{Z}_{L_r}(\alpha)] = [\Phi^{-1}(1 - \tilde{P}_{L_r}(\alpha)), \Phi^{-1}(1 - \tilde{P}_{L_l}(\alpha))] \\ &= \left[\Phi^{-1}\left(1 - \frac{\sum_{i=1}^k \tilde{p}_{L_{ir}}(\alpha)}{k}\right), \Phi^{-1}\left(1 - \frac{\sum_{i=1}^k \tilde{p}_{L_{il}}(\alpha)}{k}\right) \right],\end{aligned}\quad (6)$$

$$\begin{aligned}\tilde{Z}_U(\alpha) &= [\tilde{Z}_{U_l}(\alpha), \tilde{Z}_{U_r}(\alpha)] = [\Phi^{-1}(1 - \tilde{P}_{U_r}(\alpha)), \Phi^{-1}(1 - \tilde{P}_{U_l}(\alpha))] \\ &= \left[\Phi^{-1}\left(1 - \frac{\sum_{i=1}^k \tilde{p}_{U_{ir}}(\alpha)}{k}\right), \Phi^{-1}\left(1 - \frac{\sum_{i=1}^k \tilde{p}_{U_{il}}(\alpha)}{k}\right) \right].\end{aligned}\quad (7)$$

At the end, $\tilde{C}_{pkA}$ can be defined as what follows;

$$\tilde{C}_{pkA} = \frac{\min\{\tilde{Z}_L, \tilde{Z}_U\}}{3}, \quad (8)$$

where, $\min\{\tilde{Z}_L, \tilde{Z}_U\}$ is gain based on the ranking function. The $\tilde{C}_{pkA}$ would be at least approximately one to assess the process is capable. To make a comparison between $\tilde{C}_{pkA}$ and $\tilde{1}$, ranking function is utilized.

3.2 Fuzzy estimation of the vector $[C_{pM}, PV, LI]$

Fuzzy estimation of the first component of $[C_{pM}, PV, LI]$ is defined as the following;

$$\tilde{C}_{pM} = \left[\frac{\bigotimes_{i=1}^k (\widetilde{USL}_i \ominus \widetilde{LSL}_i)}{\prod_{i=1}^k (UPL_i - LPL_i)} \right]^{\frac{1}{k}}. \quad (9)$$

$\tilde{C}_{pM}$ is gain by placing its α -cut intervals, one on top of the other. The α -cut level sets are obtained by

$$\tilde{C}_{pM}(\alpha) = \left[\left(\frac{\prod_{i=1}^k (\widetilde{USL}_{il}(\alpha) - \widetilde{LSL}_{ir}(\alpha))}{\prod_{i=1}^k (UPL_i - LPL_i)} \right)^{\frac{1}{k}}, \left(\frac{\prod_{i=1}^k (\widetilde{USL}_{ir}(\alpha) - \widetilde{LSL}_{il}(\alpha))}{\prod_{i=1}^k (UPL_i - LPL_i)} \right)^{\frac{1}{k}} \right]. \quad (10)$$

LPL_i and UPL_i are as defined in the subsection 2.2. If the process spread is more than the tolerance spread, $\tilde{C}_{pM}$ would be less than $\tilde{1}$.

For the second component, we test the hypothesis $H_0 : \tilde{\mu} = \tilde{\mu}_0$ versus $H_1 : \tilde{\mu} \neq \tilde{\mu}_0$, where $\tilde{\mu}_0 = [\tilde{\mu}_{01}, \tilde{\mu}_{02}, \dots, \tilde{\mu}_{0k}]'$, and $\tilde{\mu}_{0i} = (\widetilde{LSL}_i \oplus \widetilde{USL}_i)/2$. To obtain $\widetilde{PV}$, first we get its α -cut level sets as what follows and then, place these intervals, one on top of the other, that has the form $\widetilde{PV}(\alpha) = [\widetilde{PV}_l(\alpha), \widetilde{PV}_r(\alpha)]$, where,

$$\begin{aligned}\widetilde{PV}_l(\alpha) &= \min \left\{ p \left(F_{k, n-k} > \frac{n-k}{k(n-1)} \begin{pmatrix} \bar{Y}_1 - m_1 \\ \bar{Y}_2 - m_2 \\ \vdots \\ \bar{Y}_k - m_k \end{pmatrix} \Sigma_{\bar{Y}}^{-1} \begin{pmatrix} \bar{Y}_1 - m_1 \\ \bar{Y}_2 - m_2 \\ \vdots \\ \bar{Y}_k - m_k \end{pmatrix} \right); \right. \\ &\quad \left. m_i \in \tilde{\mu}_{0ij}(\alpha), \quad i = \{1, 2, \dots, k\}, \quad j = \{l, r\} \right\},\end{aligned}\quad (11)$$

$$\begin{aligned} \widetilde{PV}_r(\alpha) = & \max \left\{ p \left(F_{k,n-k} > \frac{n-k}{k(n-1)} \begin{pmatrix} \bar{Y}_1 - m_1 \\ \bar{Y}_2 - m_2 \\ \vdots \\ \bar{Y}_k - m_k \end{pmatrix}' \Sigma_{\bar{Y}}^{-1} \begin{pmatrix} \bar{Y}_1 - m_1 \\ \bar{Y}_2 - m_2 \\ \vdots \\ \bar{Y}_k - m_k \end{pmatrix} \right); \right. \\ & \left. m_i \in \tilde{\mu}_{0ij}(\alpha), \quad i = \{1, 2, \dots, k\}, \quad j = \{l, r\} \right\}. \end{aligned} \quad (12)$$

Several answers are obtained based on above equations, but we choose the ones satisfy in the following condition $\widetilde{PV}_l(\beta) \leq \widetilde{PV}_l(\lambda) \leq \widetilde{PV}_l(1) = \widetilde{PV}_r(1) \leq \widetilde{PV}_r(\lambda) \leq \widetilde{PV}_r(\beta)$, for all $\beta, \lambda \in [0, 1]$, $\beta < \lambda$.

To make a decision, first we detect a degree of uncertainty such $\gamma \in (0, 1)$. Therefore, we have a three decision problem as follows;

- If $\widetilde{PV}_l(\gamma) > 0.05$, then the process mean vector is not far from the target vector.
- If $\widetilde{PV}_r(\gamma) < 0.05$, then the process mean vector is far from the target vector.
- If $\widetilde{PV}_l(\gamma) \leq 0.05 \leq \widetilde{PV}_r(\gamma)$, then we can not come to a clear decision. In such cases, one may take more samples and follow the procedure until making a decision.

For the third component, we defuzzify specification limits. To do this, we utilize the ranking function. Overall, the process supposed to be capable if $\tilde{C}_{pM} \geq_R \tilde{1}$, $\widetilde{PV}_l(\gamma) > 0.05$ and $LI = 1$.

4 Numerical example

In this section, we apply our proposed indices to asses the capability of a process discussed by (6) for monitoring simple linear profile $3 + 2x_i + \epsilon_{ij}$, where $\epsilon_{ij} \sim N(0, 1)$ and four levels are considered for independent variable. Fuzzy specification limits of the response variable at each level of the independent variable are represented in Table 1.

Table 1: *Fuzzy specification limits for each level of independent variable.*

i	X_i	$\widetilde{LSL}$	$\widetilde{USL}$
1	2	$T(1.5, 2.5, 3.5)$	$T(9, 10, 11)$
2	4	$T(5.85, 6.85, 7.85)$	$T(13.35, 14.35, 15.35)$
3	6	$T(10.25, 11.25, 12.25)$	$T(17.75, 18.75, 19.75)$
4	8	$T(15.25, 16.25, 17.25)$	$T(22.75, 23.75, 24.75)$

A random sample of size 20 is gathered and the response variables are provided in Table 2. We obtain $\tilde{\hat{C}}_{pkA} = T(0.622, 0.844, 1.074)$, i.e., “approximately 0.844”. Figure 1 shows these membership functions of $\tilde{P}_L$, $\tilde{P}_L$ and $\tilde{\hat{C}}_{pkA}$. Based on ranking function, we get $\tilde{\hat{C}}_{pkA} <_R \tilde{1}$ and conclude that the process is non-capable.

Furthermore, the fuzzy capability vector is gain as $[Approximately\ 0.825, Approximately\ 0.316, 0]$. The membership function plots of $\tilde{C}_{pM}$ and $\widetilde{PV}$ is depicted in Figure 2. Applying the ranking function for the first component, it is concluded that $\tilde{C}_{pM} <_R \tilde{1}$. Therefore, the process spread is more than the tolerance spread. For the second component, a certain degree of imprecision on sample data is set to $\gamma = 0.85$. Hence, the 0.85-level set of $\widetilde{PV}$ is $\widetilde{PV}(0.85) = [0.22785, 0.42423]$. Since $\widetilde{PV}_l(0.85) > 0.05$, it is concluded that the process mean vector is not far from the target vector. As a whole, it is concluded that the proecess is non-capable and the source of incapability is an increase in the process dispersion.

Table 2: *Response variable values for each level of independent variable for the process.*

j	y_{1j}	y_{2j}	y_{3j}	y_{4j}
1	8.19837	10.47989	16.10724	18.09119
2	5.41900	13.50778	13.33350	20.76592
3	6.69395	12.61018	14.70334	17.84424
4	6.52728	11.72180	12.04253	19.58016
5	6.18652	11.19089	15.45497	20.34619
6	7.89796	12.31044	16.98637	19.52664
7	7.52679	10.70415	12.81237	20.19214
8	6.13878	11.74836	15.21016	17.89524
9	6.56923	11.61733	14.52857	18.87686
10	6.69420	10.76137	15.59695	20.68131
11	8.16622	10.80810	13.39272	16.41702
12	6.61339	12.10850	14.76365	20.73736
13	6.88168	9.89736	14.66455	17.47383
14	5.86898	8.92178	14.52292	19.65359
15	7.64431	11.49970	16.12611	19.52600
16	6.65401	9.34401	13.44022	18.85888
17	4.53373	12.07692	13.75410	17.30811
18	8.15546	10.12848	16.08425	19.51222
19	6.64076	12.74778	14.40005	19.54112
20	7.02704	12.53329	16.88411	17.70131

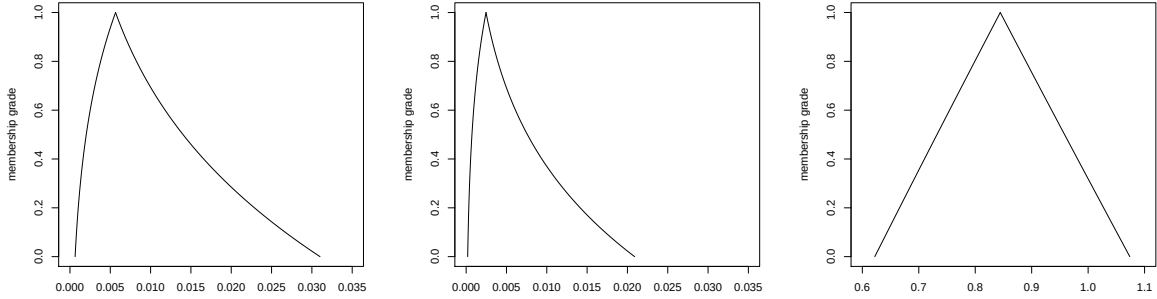


Figure 1: *The membership functions plots of $\tilde{P}_L$ (the left one), $\tilde{P}_U$ (the middle one) and $\tilde{C}_{pkA}$ (the right one).*

5 Conclusion

In this article, two fuzzy methods were developed to measure the process capability of simple linear profile for the processes that the specification limits of the response variable are fuzzy numbers. Furthermore, the new methods were applied for numerical example to demonstrate the performance of the new indices.

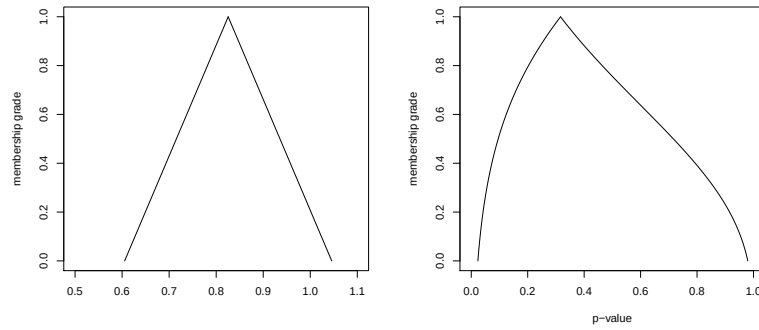


Figure 2: The membership functions plots of $\tilde{C}_{pM}$ (the left one) and $\widehat{PV}$ (the right one).

References

- [1] Draper, N.R. and Smith, H. (2011), *Applied Regression Analysis*, Third Edition, John Wiley & Sons, Inc.
- [2] Ebadi, M. and Shahriari, H. (2013), A process capability index for simple linear profile, *International Journal of Advanced Manufacturing Technology*, 64, 857-865.
- [3] Fortemps, P. and Roubens, M. (1996), *Ranking and defuzzification methods based on area compensation*, Fuzzy Sets and Systems, 82: 319-330
- [4] Golberg, M.A. and Cho, H.A. (2004), *Introduction to regression analysis*, WIT, Southampton.
- [5] Hardle, W. and Simar, L. (2007), *Applied multivariate statistical analysis*, Springer, Berlin.
- [6] Kang, L. and Albin, S.L. (2000), On-line monitoring when the process yields a linear profile, *Journal of Quality Technology*, 32, 418-426
- [7] Kunter, M.H., Nachtsheim, C.J., Neter, J. nad Li W. (2005), *Applied linear statistical models*, McGraw-Hill, Boston.
- [8] Pearn, W.L. and Kotz, S. (2006), *Encyclopedia and handbook of process capability indices. Series on quality, reliability and engineering statistics*, vol 12. World Scientific Publishing, Singapore.
- [9] Shahriari, H. and Lawrence, F.P. (1995), A multivariate process capability vector. In: *4th Industrial Engineering Research Conference*, 304-309

A recurrent neural network model to solve the fuzzy shortest path problem

Mohammad Eshaghnezhad¹, Freydoon Rahbarnia¹, Sohrab Effati¹

¹Department of Applied Mathematics, Ferdowsi University of Mashhad

Abstract

In the presence research, we are going to obtain the solution of the fuzzy shortest path (FSP) problem. According to our search in the scientific reported papers, Our focus on the paper is to give a recurrent neural network model to solve the FSP. Moreover, the proposed model is proved to be globally stable. Finally, several numerical simulations are given to show the performance of the proposed approach.

Keywords: Fuzzy shortest path problem , linear programming problem , weighting problem , recurrent neural network , globally convergent.

1 Introduction

The shortest path (SP) problem is related to finding the shortest path in a given network from a specified origin to a specified destination. Also, in this case we want to minimize the total cost. The problem arise in a wide variety of scientific and engineering applications such as traffic routing in communication networks, vehicle routing in transportation systems, economic, and etc (see (3; 1; 2)).

Let $G = (V, E)$ be a graph, where V is the set of vertices and E is the set of edges. A path between two nodes is an alternating sequence of vertices and edges starting and ending with the vertices. The distance (cost) of a path is the sum of the weights of the edges on the path. However, since there can be more than one path between any two vertices, the problem of finding a path with the minimum cost between two specified vertices makes sense. This is the so-called shortest path problem. In real world problem the arc length of the network may represent the time or cost which is not stable in the entire situation, hence it can be considered to be a fuzzy set. The fuzzy shortest path (FSP) problem was first discussed by Dubois (4) in 1980. Although the shortest path distance can be obtained, corresponding shortest path cannot be identified. Klein (7) proposed a dynamical programming recursion-based fuzzy algorithm that specified each arc length within an integer value from one to a fixed number.

¹Speaker: m_shaghnezhad@yahoo.com

Neural networks concepts were found in 1985 by Tank and Hopfield (5). There exist some neural network models for solving the quadratic optimization problem (7; 4). Motivated by the above discussions, in the present paper, we propose a new one-layer structure recurrent neural network model to solve the FSP problem. As we know, there is not an attempt for solving the FSP problem by recurrent neural network models. This methodology is going to study the recurrent neural network model to solve the FSP problem. In addition, to show the stability property of the proposed model, a new Lyapunov function is given. As we mentioned, the main advantage of the neural network approach is inherently parallel and distributed nature.

2 Fuzzy Shortes Path Problem

In this section a mathematical formulation of the FSPP is presented.

In directed network $G = (V, E)$, each arc is denoted by an ordered pair (i, j) , where $i, j \in V$. Note that, there is only one directed arc (i, j) from i to j . Also, node 1 is the source node and assume that t is the destination node. Let $\tilde{c}_{ij}$ denotes a triangular fuzzy number corresponding to the edge (i, j) , associated with the length necessary to traverse (i, j) from i to j . In real world problems, the length associates with the cost, the time, the distance, etc. Then, the FSPP is formulated as the following linear programming problem:

$$\begin{aligned} \min \quad & \tilde{f} = \sum_{(i,j) \in E} \tilde{c}_{ij} x_{ij} \\ \text{s.t.} \quad & \sum_j x_{ij} - \sum_j x_{ji} = \begin{cases} 1, & \text{if } i = 1, \\ 0, & \text{if } i \neq 1, r \quad (i = 1, \dots, r), \\ -1, & \text{if } i = r, \end{cases} \\ & x_{ij} = 0 \text{ or } 1, \quad \text{for } (i, j) \in E, \end{aligned} \tag{1}$$

where r is the number of nodes, and the symbol $\sum$ in the objective function refers to an addition $\oplus$ of fuzzy numbers. Also, x_{ij} denotes the decision variable associated with the edge from vertices i to j as defined below:

$$x_{ij} = \begin{cases} 1, & \text{if the edge from vertices } i \text{ to } j \text{ is in the path;} \\ 0, & \text{otherwise.} \end{cases}$$

Because of the total unimodality property of the constraint coefficient matrix defined in (1) (3), the integrality constraint in the shortest path problem formulation can be equivalently replaced with the non-negativity constraint, if the shortest path is unique. By the above discussion the following equivalent linear programming problem can be formulated:

$$\begin{aligned} \min \quad & \sum_{i=1}^r \sum_{\substack{j=1 \\ j \neq i}}^r \tilde{c}_{ij} x_{ij} \\ \text{s.t.} \quad & \sum_{\substack{k=1 \\ k \neq i}}^r x_{ik} - \sum_{\substack{l=2 \\ l \neq i}}^r x_{li} = \delta_{i1} - \delta_{ir}, \quad i = 1, \dots, r, \\ & x_{ij} \geq 0, \quad i \neq j, \quad i = 1, \dots, r-1, \quad j = 1, \dots, r, \end{aligned} \tag{2}$$

where $\delta_{ij} = 1$ for $i = j$ and equal to 0 for $i \neq j$. We can write the problem (2) as a interval program:

$$\begin{aligned}
\min \quad & \sum_{i=1}^r \sum_{\substack{j=1 \\ j \neq i}}^r [\underline{c}_{ij}(\alpha)x_{ij}, \bar{c}_{ij}(\alpha)x_{ij}] \\
s.t. \quad & \sum_{\substack{k=1 \\ k \neq i}}^r x_{ik} - \sum_{\substack{l=2 \\ l \neq i}}^r x_{li} = \delta_{i1} - \delta_{ir}, \quad i = 1, \dots, r, \\
& x_{ij} \geq 0, \quad 0 \leq \alpha \leq 1, \quad i \neq j, \quad i = 1, \dots, r-1, \quad j = 1, \dots, r,
\end{aligned} \tag{3}$$

By reformulating (3) into a weighting problem, we have:

$$\begin{aligned}
\min \quad & \sum_{i=1}^r \sum_{\substack{j=1 \\ j \neq i}}^r (w_1 \underline{c}_{ij}(\alpha) + w_2 \bar{c}_{ij}(\alpha)) x_{ij} \\
s.t. \quad & \sum_{\substack{k=1 \\ k \neq i}}^r x_{ik} - \sum_{\substack{l=2 \\ l \neq i}}^r x_{li} = \delta_{i1} - \delta_{ir}, \quad i = 1, \dots, r, \\
& x_{ij} \geq 0, \quad 0 \leq \alpha \leq 1, \quad i \neq j, \quad i = 1, \dots, r-1, \quad j = 1, \dots, r,
\end{aligned} \tag{4}$$

where $w_1 + w_2 = 1$. Based on the edge path representation, the dual shortest path problem can be formulated as a linear programming problem as follows:

$$\begin{aligned}
\max \quad & y_r - y_1 \\
s.t. \quad & y_j - y_i \leq w_1 \underline{c}_{ij}(\alpha) + w_2 \bar{c}_{ij}(\alpha), \quad i \neq j, \quad i, j = 1, 2, \dots, r, \\
& 0 \leq \alpha \leq 1,
\end{aligned} \tag{5}$$

where y_i denotes the dual decision variable associated with vertex i and $y_i - y_1$ is the shortest distance from vertex 1 to vertex i at optimality.

3 Recurrent Neural Network

$$\begin{aligned}
\min \quad & \tilde{c}x \\
s.t. \quad & Ax = b, \\
& x \geq 0.
\end{aligned} \tag{6}$$

Similar to the structure of the problems (3)-(4) the problem (6) can be written as follows:

$$\begin{aligned}
\min \quad & (w_1 \underline{c} + w_2 \bar{c})x \\
s.t. \quad & Ax = b, \\
& x \geq 0,
\end{aligned} \tag{7}$$

where $x, c \in \mathbb{R}^n$, $A \in \mathbb{R}^{m \times n}$, and $b \in \mathbb{R}^m$. The KKT optimality conditions for (7) can be obtained as follow:

$$\begin{aligned} & -(w_1 \underline{c} + w_2 \bar{c}) + A^T y = 0 \\ & -(w_1 \underline{c} + w_2 \bar{c}) + A^T y \geq 0 \\ & Ax = 0 \\ & x \geq 0, \end{aligned} \tag{8}$$

where y (is free) is the Lagrange multiplier. For introducing the recurrent neural network model we need some preliminaries. A mixed complementarity problem on $\mathbb{R}_+^n$ is defined as follow,

$$\begin{aligned} & F_i(x^*)x_i^* = 0, \\ & F_i(x^*)x_i^* \geq 0, \quad x_i^* \geq 0, \quad i \in I, \\ & F_j(x^*)x_j^* = 0, \quad j \in L - I, \end{aligned}$$

where, $L = \{1, \dots, n + m\}$.

Here, we are going to rewrite the KKT optimality conditions (8) as a mixed linear complementarity problem (MLCP). For this purpose, we set the vectors $F(w)$ and w as follow:

$$F(w) = \begin{bmatrix} -(w_1 \underline{c} + w_2 \bar{c}) + A^T y \\ -Ax \end{bmatrix}, \quad w = \begin{bmatrix} x \\ y \end{bmatrix}.$$

Also, it can be restated as the matrix form below:

$$F(w) = \begin{bmatrix} 0_{(m+1) \times (m+1)} & A^T \\ -A & 0 \end{bmatrix} \begin{bmatrix} x \\ y \end{bmatrix} + \begin{bmatrix} -(w_1 \underline{c} + w_2 \bar{c})^T \\ 0 \end{bmatrix}.$$

Consider the following matrix and vector:

$$M = \begin{bmatrix} 0_{(m+1) \times (m+1)} & A^T \\ -A & 0 \end{bmatrix}, \quad q = \begin{bmatrix} -(w_1 \underline{c} + w_2 \bar{c})^T \\ 0 \end{bmatrix}.$$

Therefore, $F(w)$ can be reformulated as $F(w) = Mw + q$ where M is a positive semidefinite matrix. Thus, we can restated the KKT optimality conditions as the following MLCP:

$$\begin{cases} F(w) = Mw + q, \\ F_i w_i = 0, & F_i \geq 0, \quad w_i \geq 0, \quad i = 1, 2, \dots, m, \\ F_j = 0, & j = m + 1, \dots, m + n. \end{cases} \tag{9}$$

Now, we define the function NCP function is the min function, which is defined as:

$$\Psi(a, b) = \min(a, b).$$

Also, the $\Psi_{\min}(a, b)$ can be shown with different representation as follow:

$$\Psi_{\min}(a, b) = \min\{a, b\} = (a + b - |a - b|)/2$$

Now we describe a recurrent neural network for solving (MLCP) (22) whose state variable is defined by the following system: Here, for solving the MLCP (9), the following dynamical model is given:

$$\frac{dw}{dt} = -\eta(I + M^T)H(w), \quad (10)$$

where $w = (x, z), \eta > 0$ and F is defined as follows:

$$H(w) = \phi_{min}(w_i, F_i(w)) \quad i = 1, 2, \dots, m \quad (11)$$

$$H_i(w) = F_i(w) \quad i = m + 1, m + 2, \dots, m + n, \quad (12)$$

4 Convergence Analysis

Here, we study the stability analysis of the proposed model (10). For introducing the recurrent neural network model we need some preliminaries. Here, we study the stability analysis of the proposed model (10). x^* is the KKT point of (8), if and only if, $w^* = (x^*, y^*)^T$ satisfies the equation $H(w^*) = 0$. Right hand side of (10) is Lipschitzian. $w^* = (x^*, y^*)^T$ is the equilibrium point of the recurrent neural network model (10), if and only if, satisfies the KKT conditions (8)

There is a unique Lipschitz continuous solution $x(t)$ for the neural network model (10). The model (10) for $0 < \eta$ is stable in the sense of Lyapunov and globally convergent.

5 Numerical Example

In this section, we present an illustrative example in order to demonstrate the applicability and the performance of the recurrent neural network model. Let us consider an acyclic network that possesses an arc with a negative cost in networks with such a characteristic as shown in Fig. 1. Also, the arc lengths are displayed in Table 1. We apply the proposed neural network model in (9) to solve this

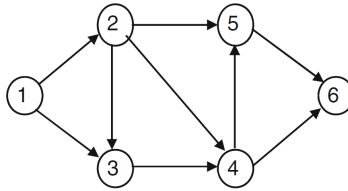


Figure 1: *Acyclic network of Example 5*

program. Fig. 2 displays the transient behaviour of $x(t)$ based on the neural network model (9) with random initial point, $\alpha = 1$, $w_1 = w_2 = \frac{1}{2}$, and $\gamma = 2$. It means the optimal solution of the problem for $\alpha = 1$ (when the fuzziness goes to zero), choosing the parameters $w_1 = w_2 = \frac{1}{2}$, and $\gamma = 2$ is shown in Fig 2. Moreover, Fig. 1 displays the transient behaviour of $x(t)$ based on the neural network model (9) with random initial point, $\alpha = 0.5$, $w_1 = w_2 = \frac{1}{2}$, and $\gamma = 2$. Also, we give the solutions for various values of α and setting parameters $w_1 = w_2 = \frac{1}{2}$ and $\gamma = 2$ in Table 2. According to the Table 2, the α -cut of f and x_{ij} represent the possibility that the value will appear in the associated range.

Table 1: *Fuzzy arc lengths of acyclic network in Example 5*

Source node	Destination node	Arc cost
1	2	(1, 2, 3)
1	3	(5, 7, 9)
2	3	(1, 4, 9)
2	4	(10, 11, 12)
2	5	(5, 6, 7)
3	4	(8, 9, 10)
4	5	(-9, -8, -7)
4	6	(11, 13, 14)
5	6	(8, 9, 10)

Table 2: *The solution of the neural network model (9) for different values of α , $w_1 = w_2 = \frac{1}{2}$, and $\gamma = 2$ in Example 5.*

$x_{ij} \backslash \alpha$	0	0.1	0.3	0.5	0.7	0.9	1
x_{12}	0.9996	0.9994	0.9991	0.9989	0.9988	0.9986	0.9986
x_{13}	0	0	0	0	0	0	0
x_{23}	0	0	0	0	0	0	0
x_{24}	0.9995	0.9994	0.9990	0.9989	0.9985	0.9985	0.9983
x_{25}	0	0	0	0	0	0	0
x_{34}	0	0	0	0	0	0	0
x_{45}	1.0002	0.0003	1.0007	1.0009	1.0009	1.0010	1.0011
x_{46}	0	0	0	0	0	0	0
x_{56}	1.0006	1.0006	1.0008	1.0011	1.0012	1.0014	1.0014

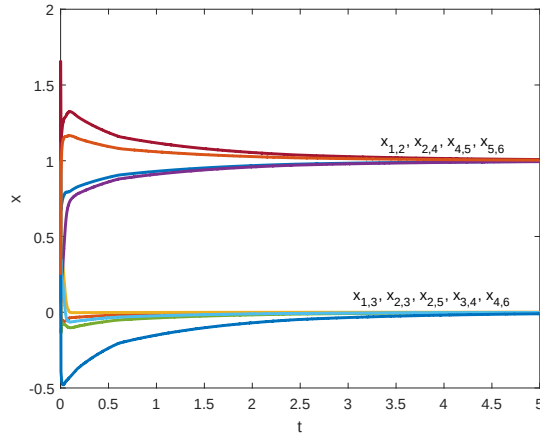


Figure 2: *Transient behaviours of the neural network (9) with random initial points, $\alpha = 1$, $w_1 = w_2 = \frac{1}{2}$, and $\gamma = 3$ in Example 5*

References

- [1] R. K. Ahuja, T. L. Magnanti, and J. B. Orlin, *Network Flows: Theory, Algorithms, and Applications*, Prentice Hall, Englewood Cliffs, New Jersey, 1993.

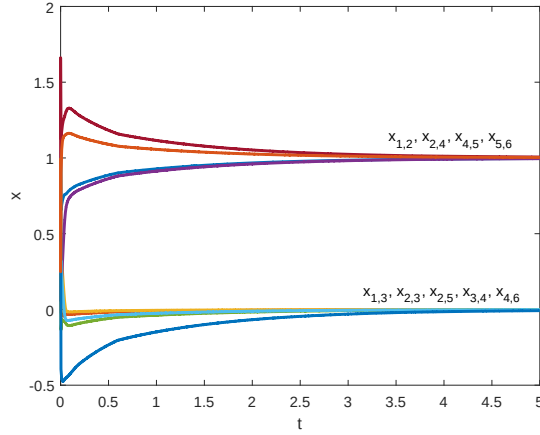


Figure 3: *Transient behaviours of the neural network (9) with random initial points, $\alpha = \frac{1}{2}$, $w_1 = w_2 = \frac{1}{2}$, and $\gamma = 3$ in Example 5*

- [2] J. K. Antonio, G. M. Huang, and W. K. Tsai, A fast distributed shortest path algorithm for a class of hierarchically clustered data networks, *IEEE Transactions on Computers* **41** (6) (1992), 710-724.
- [3] M. S. Bazaraa, C. Shetty, and H. D. Sherali, *Nonlinear Programming, Theory and Algorithms*, New York: John Wiley and Sons, 1979.
- [4] D. Dubois and H. Prade, *Fuzzy Sets and Systems: Theory and Applications*, New York: Academic Press, 1980.
- [5] M. Eshaghnezhad, S. Effati, and A. Mansoori, A Neurodynamic Model to Solve Nonlinear Pseudo-Monotone Projection Equation and Its Applications, *IEEE Transactions on Cybernetics* **47** (10) (2016), 3050-3062.
- [6] J. J. Hopfield and D. W. Tank, Neural computation of decisions in optimization problems, *Biological Cybernetics* **52** (1985), 141-152.
- [7] C. M. Klein, Fuzzy shortest paths, *Fuzzy Sets and Systems* **39** (1991), 27-41.
- [8] A. Mansoori, S. Effati, and M. Eshaghnezhad, A neural network to solve quadratic programming problems with fuzzy parameters, *Fuzzy Optimization and Decision Making* **17** (1) (2018), 75-101.

JB estimation approach applied to fuzzy regression modeling

M. Kashani¹ ¹, M. Arashi ², M.R. Rabiei³

¹Department of statistics, Faculty of Mathematical Sciences, Shahrood University of Technology, Shahrood, Iran

²Department of statistics, Faculty of Mathematical Sciences, Shahrood University of Technology, Shahrood, Iran

³Department of statistics, Faculty of Mathematical Sciences, Shahrood University of Technology, Shahrood, Iran

Abstract

When the sample size is small, optimal estimation in the context of regression modeling is important. In a typical fuzzy regression with crisp input and fuzzy error term, we implement the fuzzy Jackknife-after-Bootstrap (JB) technique to estimate the fuzzy parameters. Numerical example illustrate the proposed JB approach improves the estimates of fuzzy parameters.

Keywords: Fuzzy regression, Fuzzy least-square method, Jackknife-after-Bootstrap (JB) technique method, α -cut.

1 Introduction

In 1982, Tanaka and his colleagues proposed a mathematical relation for fuzzy regression modeling (FRM) afterward many efforts have been done to refine their FRM mechanism, see (9) for an extensive review. However, when the size of samples is small, resampling technique must be used in estimating the parameters of a fuzzy regression model. Resampling techniques such as bootstrap were applied to FRM by (1), (8), and (7). Recently (6) proposed the fuzzy Jackknife-after-Bootstrap (JB) technique to enhance the efficiency of bootstrap method.

The structure of the paper is as follows: In Section 2, some basic concepts of fuzzy theories are briefly reviewed. In Section 3, the fuzzy analogs of the normal equations of the least squares method are addressed and the fuzzy least squares estimators are obtained. Further the JB method is applied to estimation of parameters. In Section 4, numerical examples are presented for illustration purpose of the JB technique and providing optimal model criteria in comparison with bootstrap method for each α -cut. Finally, Section 5 summarizes the main results and draws conclusions.

¹Speaker: email@email.ac.ir

2 Preliminaries

In this section, we introduce some definitions regarding the fuzzy sets and the fuzzy numbers, as well as some basic concepts of fuzzy theory.

2.1 Resampling methods

Resampling methods are robust procedures to increase accuracy in model estimation with low sample sizes. The jackknife and the bootstrap methods assume that the quadrats are independent and represent a sample from the same distribution. No assumption, however, is made about the relationships among species within a quadrat.

Suppose one is interested in estimating some parameters, θ , using $\hat{\theta} = F(X_1, X_2, \dots, X_n)$, where $\mathbf{X} = (X_1, X_2, \dots, X_n)$ is a sample of n independent random variables with cumulative distribution function $F(\mathbf{X}, \theta)$. Assume that $\hat{\theta}$ is a reasonably good estimate of θ . To obtain the jackknife estimate using the observations $\mathbf{x} = (x_1, x_2, \dots, x_n)$, one performs the following algorithm (12).

- (i) Remove one of the observations, say, x_i .
- (ii) Compute the estimate of θ based on $(x_1, x_2, \dots, x_{i-1}, x_{i+1}, \dots, x_n)$, and denote it by $\hat{\theta}_{-i}$.
- (iii) Compute the pseudo value $\hat{\theta}_i = n\hat{\theta} - (n-1)\hat{\theta}_{-i}$.
- (iv) Repeat (i) to (iii) n times for $i = 1, 2, \dots, n$.

The jackknife estimate $J_n^1(\theta)$, is then given by $J_n^1(\theta) = \frac{1}{n} \sum_{i=1}^n \hat{\theta}_i$. Bootstrap is developed as a method related to the jackknife but more widely applicable and dependable. The bootstrap method typically requires simulation methods for estimation of the parameter and its variance. The bootstrap procedure will be based on (7).

2.1.1 Jackknife-after-Bootstrap method

Let us denote $\text{se}(\hat{\theta})$, the sample standard deviation of B bootstrap replicates of $\hat{\theta}$. Now, if we leave out the i^{th} observation, the algorithm for estimation of standard errors is to resample, B replicates from $n-1$ remaining observations (for each i). In fact, the importance of this method of obtaining an approximation of the standard error is related to the bootstrap samples. The jackknife-after-bootstrap (JB) computes an estimate for each “leave-one-out” sample. This method is due to (5). Let $J(i)$ denote the indices of bootstrap samples that do not contain x_i , and let $B(i)$ denote the number of bootstrap samples that do not contain x_i . Then we can compute the jackknife replication leaving out the $B - B(i)$ samples that contain x_i .

2.2 Fuzzy definition

1.2. A fuzzy subset $\tilde{A}$ with membership function $\xi_{\tilde{A}} : \mathbb{R}^n \rightarrow [0, 1]$ is called a fuzzy number, if $\xi_{\tilde{A}}$ satisfies

- (i) $\tilde{A}$ is normal,
- (ii) $\tilde{A}$ is convex,
- (iii) $\xi_{\tilde{A}}$ is semicontinuous,
- (iv) $\text{supp}(\tilde{A}) := \{x \in \mathbb{R}^n : \xi_{\tilde{A}}(x) > 0\}$ is compact.

2.2. The α -cut of a fuzzy number $\tilde{A}$ is a non-fuzzy set defined as

$$A_\alpha := \{x \in \mathbb{R}^n \mid \xi_{\tilde{A}}(x) \geq \alpha, 0 \leq \alpha \leq 1\}.$$

The set of all fuzzy numbers will be denoted by $\mathcal{F}_c(\mathbb{R})$. The α -cut A_α of $\tilde{A} \in \mathcal{F}_c(\mathbb{R})$ is a closed and bounded interval for each $\alpha \in [0, 1]$. Hence a fuzzy number $\tilde{A} \in \mathcal{F}_c(\mathbb{R})$ is completely determined by the end points of the intervals $A_\alpha = [A_\alpha^L, A_\alpha^U]$. The arithmetic operations for two fuzzy numbers A_α and B_α are defined in the standard way, in terms of the α -cuts for $\alpha \in [0, 1]$.

Addition: $A_\alpha + B_\alpha = [A_\alpha^L + B_\alpha^L, A_\alpha^U + B_\alpha^U]$.

Scalar multiplication: for given $k \in \mathbb{R}$,

$$k \cdot A_\alpha = \begin{cases} [kA_\alpha^L, kA_\alpha^U] & \text{if } k \geq 0 \\ [kA_\alpha^U, kA_\alpha^L] & \text{if } k < 0 \end{cases}$$

Scalar addition: $k + A_\alpha = [k + A_\alpha^L, k + A_\alpha^U]$.

Using the general Hukuhara difference (13), the subtraction of interval is defined as follows:

$$[A_\alpha^L, A_\alpha^U] \ominus_g [B_\alpha^L, B_\alpha^U] = [C_\alpha^-, C_\alpha^+]$$

where $C_\alpha^- = \min \{A_\alpha^L - B_\alpha^L, A_\alpha^U - B_\alpha^U\}$, $C_\alpha^+ = \max \{A_\alpha^L - B_\alpha^L, A_\alpha^U - B_\alpha^U\}$. An alternative representation of a closed interval $A_\alpha = [A_\alpha^L, A_\alpha^U]$ is by using the midpoint and the width, i.e.,

$$A_\alpha^C = \frac{1}{2} (A_\alpha^L + A_\alpha^U) \quad \text{and} \quad A_\alpha^W = \frac{1}{2} (A_\alpha^U - A_\alpha^L).$$

Therefore, the interval A_α can be equivalently expressed as the vector $A_\alpha = (A_\alpha^C, A_\alpha^W)$, $A_\alpha^W \geq 0$. In this paper, the distance between α -cuts of two fuzzy numbers $\tilde{A}$ and $\tilde{B}$ as Euclidean distance is applied by

$$d(A_\alpha, B_\alpha) = \sqrt{\{A_\alpha^C - B_\alpha^C\}^2 + \{A_\alpha^W - B_\alpha^W\}^2}.$$

3 FRM

In this paper, we will use the fuzzy regression model as follows:

$$\tilde{y}_i = \tilde{\beta}_0 + \tilde{\beta}_1 x_i + \tilde{\varepsilon}_i, \quad i = 1, \dots, n \quad \text{and} \quad \alpha \in [0, 1],$$

where $\tilde{y}_i$ is the fuzzy output and x_i is a non-fuzzy input variables. Also $\tilde{\beta}_0$ and $\tilde{\beta}_1$ are the fuzzy coefficient and $\tilde{\varepsilon}_i$ is an error without assumption of distribution.

3.1 Estimation of fuzzy parameters

Based on the least square method, distance defined above and the fact that α -cuts generate non-fuzzy intervals, now we have:

$$y_{i\alpha} = \beta_{0\alpha} + \beta_{1\alpha} x_i + \varepsilon_{i\alpha}, \quad \hat{\beta}_{0\alpha}^K = \bar{Y}_\alpha^K - \hat{\beta}_{1\alpha}^K \bar{X}, \quad \hat{\beta}_{1\alpha}^K = \frac{\sum_{i=1}^n (x_i - \bar{x})(y_{i\alpha}^K - \bar{y}_\alpha^K)}{\sum_{i=1}^n (x_i - \bar{x})^2}$$

Where $\bar{x}$ is the mean of input values and $\bar{y}^k$ is the mean of the observed responses under the α and k were $K = L, U$, (L =lower and U =upper) so we denote the estimated closed interval $\hat{\beta}_{j\alpha}^K$ for $j = 0, 1$ by $\hat{\beta}_{j\alpha}^K = [\hat{\beta}_{j\alpha}^-, \hat{\beta}_{j\alpha}^+]$, where $\hat{\beta}_{j\alpha}^- = \min\{\hat{\beta}_{j\alpha}^K\}$ and $\hat{\beta}_{j\alpha}^+ = \max\{\hat{\beta}_{j\alpha}^K\}$. Then using these initial coefficients, we obtain bootstrap coefficients for each α -cut applying the method of (7). We will be also computing the estimates using the JB method.

3.2 GOF measures

In this paper, some measures are used for evaluating the performance of a fuzzy regression model. Let, $\tilde{Y}_i = (l_{Y_i}, m_{Y_i}, r_{Y_i})_T$ and $\hat{\tilde{Y}}_i = (l_{\hat{Y}_i}, m_{\hat{Y}_i}, r_{\hat{Y}_i})_T$ are triangular fuzzy numbers and $\tilde{Y}_i, \hat{\tilde{Y}}_i$ respectively value of i^{th} observation and estimation. Following (11), the error in estimation between the fuzzy output and its corresponding fuzzy estimated value is defined as follows

$$D_{p,q} = \frac{1}{n} \sum_{i=1}^n D_{p,q}^2(i), \quad i = 1, 2, \dots, n$$

For the triangular fuzzy number ($p = 2, q = \frac{1}{2}$), the above equation will be presented as follows.

$$D_{2,1/2}(\hat{\tilde{Y}}_i, \tilde{Y}_i) = \frac{1}{6} \left[(l_{\hat{Y}_i} - l_{Y_i})^2 + 2(m_{\hat{Y}_i} - m_{Y_i})^2 + (r_{\hat{Y}_i} - r_{Y_i})^2 + (l_{\hat{Y}_i} - l_{Y_i})(m_{\hat{Y}_i} - m_{Y_i}) + (r_{\hat{Y}_i} - r_{Y_i})(m_{\hat{Y}_i} - m_{Y_i}) \right].$$

In addition, we use the other criteria used to evaluate the method from the following meter (3).

$D^2(\hat{\tilde{Y}}_i, \tilde{Y}_i) = (m_{\hat{Y}_i} - m_{Y_i})^2 + ((m_{\hat{Y}_i} - l_{\hat{Y}_i}) - (m_{Y_i} - l_{Y_i}))^2 + ((m_{\hat{Y}_i} - r_{\hat{Y}_i}) - (m_{Y_i} - r_{Y_i}))^2$, where $D(\hat{\tilde{Y}}_i, \tilde{Y}_i)$ is a distance between the estimated and observed membership values of the response variable, and $\tilde{Y}$ is average of observations of $\tilde{Y}_i$. As some measures for investigating the GOF of the model, we suggest to use

$$SPE = \sum_{i=1}^n D^2(\hat{\tilde{Y}}_i, \tilde{Y}_i), \quad D_\alpha^2 = \sum_{i=1}^n D_{2,1/2}(\hat{\tilde{Y}}_i, \tilde{Y}_i)$$

These criteria tell us how closely a fuzzy response is related to its fuzzy estimated value.

4 Numerical Illustrations

In this section, we provide some numerical illustrations to compare the performance of the proposed JB method with the existing ones in the literature. For our purpose, we use two real data sets.

1.4. Table 1 shows the data for our FRM, which is adopted from (7).

Table 1: Cheese tasting data.

<i>Obs.</i>	x_i	$\tilde{y}_i$	<i>Obs.</i>	x_i	$\tilde{y}_i$
1	1	(6.2,6.2,6.8)	9	9	(6.0,6.2,7.1)
2	2	(6.1,6.4,6.8)	10	10	(6.4,6.7,7.0)
3	3	(5.7,5.8,6.3)	11	11	(6.6,6.9,7.2)
4	4	(6.7,6.8,7.0)	12	12	(6.7,7.0,7.3)
5	5	(5.5,5.9,6.2)	13	13	(6.3,6.5,6.9)
6	6	(5.9,6.2,6.5)	14	14	(6.5,6.8,7.3)
7	7	(6.1,6.4,6.7)	15	15	(6.1,6.3,7.4)
8	8	(6.0,6.6,6.9)			

Table 1 includes the estimates of regression coefficient. Now, we compare the results of the two methods according to the GOF criteria in Table 2.

Table 2: Model coefficients from example 1.

bootstrap technique					JB technique				
$\hat{\beta}_i^L$		α -cut	$\hat{\beta}_i^U$		$\hat{\beta}_i^L$		α -cut	$\hat{\beta}_i^U$	
$\hat{\beta}_0$	$\hat{\beta}_1$		$\hat{\beta}_0$	$\hat{\beta}_1$	$\hat{\beta}_0$	$\hat{\beta}_1$		$\hat{\beta}_0$	$\hat{\beta}_1$
5.8686	0.0330	0	6.4636	0.0490	5.8702	0.0328	0	6.5182	0.0490
5.9041	0.0330	0.1	6.4376	0.0563	5.9026	0.0331	0.1	6.4374	0.0562
5.9300	0.0339	0.2	6.3965	0.0556	5.9266	0.0343	0.2	6.3986	0.0549
5.9525	0.0353	0.3	6.3845	0.0516	5.9562	0.0350	0.3	6.3853	0.0511
5.9851	0.0353	0.4	6.3237	0.0527	5.9857	0.0352	0.4	6.3246	0.0527
6.0197	0.0356	0.5	6.3032	0.0490	6.0235	0.0353	0.5	6.3031	0.0492
6.0497	0.0357	0.6	6.2586	0.0488	6.0476	0.0360	0.6	6.2580	0.0490
6.0547	0.0361	0.7	6.2268	0.0465	6.0647	0.0376	0.7	6.2260	0.0465
6.0679	0.0374	0.8	6.1931	0.0458	6.0779	0.0394	0.8	6.1926	0.0459
6.1013	0.0381	0.9	6.1590	0.0432	6.1213	0.0391	0.9	6.1577	0.0434
6.1333	0.0411	1	6.1333	0.0411	6.1350	0.0411	1	6.1350	0.0411

From Table 2, the JB method optimizes to coefficient more, compared to the bootstrap.

Table 3: *Comparison of GOFs*

method	GOF criteria	$\alpha - cut$										
		0	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9	1.0
bootstrap	D_α^2	1.1672	1.1655	1.1685	1.1756	1.1851	1.1983	1.2163	1.2418	1.2754	1.3154	1.3239
	SPE_α	1.7222	1.7091	1.7007	1.6965	1.6947	1.6964	1.7031	1.7173	1.7397	1.7685	1.7650
JB	D_α^2	1.1666	1.1654	1.1679	1.1749	1.1849	1.1981	1.2162	1.2374	1.2633	1.2915	1.3237
	SPE_α	1.7220	1.7091	1.7003	1.6959	1.6944	1.6962	1.7030	1.7129	1.7274	1.7441	1.7220

2.4. Here, this example is from (2), when the data are provided in 4.

Table 4: *Cheese tasting data.*

Obs.	x_i	$\tilde{y}_i$	Obs.	x_i	$\tilde{y}_i$
1	1	(16.8,21,23.1)	5	5	(10.8,12,13.2)
2	2	(12.75,15,17.25)	6	6	(14.4,18,19.8)
3	3	(13.5,15,17.25)	7	7	(5.4,6,7.2)
4	4	(7.65,9,10.35)	8	8	(10.2,12,14.4)

Table 6 includes the estimates of regression coefficient. Now, we compare the results of the two methods according to the GOF criteria in Table 5.

Table 5: *Model coefficients from example 1.*

bootstrap technique					JB technique				
$\hat{\beta}_i^L$		$\alpha\text{-cut}$	$\hat{\beta}_i^U$		$\hat{\beta}_i^L$		$\alpha\text{-cut}$	$\hat{\beta}_i^U$	
$\hat{\beta}_0$	$\hat{\beta}_1$		$\hat{\beta}_0$	$\hat{\beta}_1$	$\hat{\beta}_0$	$\hat{\beta}_1$		$\hat{\beta}_0$	$\hat{\beta}_1$
14.7716	-1.1789	0	21.3847	-0.9428	14.9803	-1.2120	0	21.3989	-0.9484
14.9290	-1.1258	0.1	21.2341	-0.9727	14.9004	-1.1105	0.1	21.3187	-0.9889
15.9704	-1.2101	0.2	20.7728	-0.9653	15.9272	-1.2040	0.2	20.5919	-0.9318
15.9804	-1.1555	0.3	20.5356	-0.9697	16.4523	-1.1625	0.3	20.4348	-0.9863
16.4383	-1.1615	0.4	20.4228	-0.9896	16.4775	-1.1655	0.4	20.3652	-0.9846
16.6246	-1.1196	0.5	20.1901	-1.0437	16.6049	-1.1061	0.5	20.3049	-1.0620
17.0178	-1.1546	0.6	19.8448	-1.0631	17.0114	-1.1505	0.6	19.8450	-1.0707
17.2162	-1.1085	0.7	19.7756	-1.0455	17.2104	-0.9750	0.7	19.7002	-1.0231
17.5643	-1.1598	0.8	19.5356	-1.0969	17.4773	-1.0823	0.8	19.5399	-1.1559
18.3509	-1.1653	0.9	18.6240	-1.0799	18.4163	-1.1647	0.9	18.5423	-1.0664
18.6279	-1.1547	1	18.6279	-1.1547	18.6196	-1.1470	1	18.6196	-1.1470

From Table 5, the JB method optimizes to coefficient more, compared to the bootstrap.

Table 6: *Comparison of GOFs*

method	GOF criteria	$\alpha - cut$										
		0	0.1	0.2	0.3	0.4	0.5	0.6	0.7	0.8	0.9	1.0
bootstrap	D_{α}^2	107.1859	109.4350	108.3079	106.2875	106.2833	106.1602	105.3412	105.4772	106.3429	106.7345	106.2853
	SPE_{α}	125.0503	127.8402	127.1157	125.9311	125.7791	125.5599	124.9784	124.8682	125.1558	125.1643	124.7699
JB	D_{α}^2	107.1578	109.1664	108.0725	106.2522	106.2191	106.0002	105.2922	105.3520	105.9191	105.8283	106.2526
	SPE_{α}	125.0174	127.6860	126.9762	125.9173	125.7190	125.4495	124.9279	124.7818	124.9314	124.6882	124.7328

5 Conclusion

In this paper, we proposed a Jackknife-after-Bootstrap estimation technique, following (6), to important upon the bootstrap method. Evaluating the GOF criteria revealed the proposed technique improves estimation in the FRM.

References

- [1] Akbari MG, Mohammadalizadeh R, Rezaei M (2012), Bootstrap statistical inference about the regression coefficients based on fuzzy data, *Int. J. Fuzzy Syst.* 14, 549-556.
- [2] Chen S.P. and Dang J.F. (2008), A variable spread fuzzy linear regression model with higher explanatory power and forecasting accuracy, *Inf. Sci.* 178, 3973-3988.
- [3] Diamond P. (1988), Fuzzy least square, *Inf. Sci.* 46, 141-157.
- [4] Efrone B. (1979), Bootstrap methods: another look at the Jackknife, *Analysis Statistics*, 101-118.
- [5] Efron B. and TibshIrani R.J. (1993), *An Introduction to the Bootstrap*, Chapman & Hall/CRC, Boca Raton, FL.
- [6] Kashani M., Arashi M., Rabiei M.R. (2017), Jackknife-after-Bootstrap in Fuzzy Regression Modeling, *17th conference on fuzzy systems and 15th conference on intelligent systems*.
- [7] Lee W.J., Jung H.Y, Yoon J.H., Choi S.H. (2015), The statistical inferences of fuzzy regression based on bootstrap techniques, *Soft Comp.* 19, 883-890.
- [8] Maria B.F., Renato C., Gil G.R. (2013), Bootstrap confidence intervals for the parameters of a linear regression model with fuzzy random variables, *Towards Adv Data Anal Comb Soft Comput Stat* 285, 33-42.
- [9] Moghaddas M. (2014), Comparison of several optimization methods in fuzzy linear regression, MSc Thesis, Ferdowsi university of Mashhad, Faculty of Mathematical Sciences, Department of statistics.
- [10] Puri M.D. and Ralescu D. (1986), Fuzzy random variables, *J. Math Anal. Appl.* 114, 409-422.
- [11] Sadeghpour Gildeh, B. and Gien, D. (2002), A goodness of fit index to reliability analysis in fuzzy model, *In 3rd WSEAS international conference on fuzzy sets and fuzzy systems*, Iterlaken, Switzerland, 11-14.
- [12] Smith E.P. and Belle G. (1984), Nonparametric estimation of species richness, *Biometrics* 40,1, 119-129.
- [13] Stefanini L (2010), A generalization of Hukuhara difference for interval and fuzzy arithmetic, *Fuzzy Sets Syst.* 161, 1564-1576.
- [14] Tanaka H, Uejima S, Asai K (1982), Linear regression analysis with fuzzy model, *IEEE Syst Trans Syst Man Cybern SMC-2* 903-907.

Topic Modeling Based Plagiarism Detection in Persian Documents

Banafshe Mohammadian¹ , Jafar Khakpour², Mir Mohsen Pedram^{1 3}

^{1,2,3}Department of Electrical & Computer Engineering, Faculty of Engineering, Kharazmi University, Tehran/Iran

Abstract

Plagiarism is a rapidly growing problem in the advent of internet and information explosion. Automatic techniques for detecting plagiarized documents are crucially important. There are many machine learning algorithms suggested to automatically detect plagiarized documents in a large corpus of documents. This research aims to use machine learning algorithms such as vector space model (VSM) and singular value decomposition (SVD) algorithms for this purpose. This paper uses vector space model to represent corpus of text data, and SVD is used to analyze latent semantics of these documents. Natural language processing techniques are also used to process and clean documents, and documents are clustered by topic modeling technique to enhance its accuracy. Experimental results show the effectiveness of the proposed method.

Keywords:Plagiarism Detection, Topic Modeling, Latent Dirichlet Allocation, Singular Value Decomposition

1 Introduction

Academic plagiarism is known as using other researchers' publications and ideas to publish a new document without referring to base works or representing these works as someone else's original work (14). With the increasing volume of online data, plagiarism has become an issue on academic publishing. The large volume of text information available online makes it easy to publish plagiarized documents. This availability and abundance also makes detecting plagiarism a hard task. Information retrieval (IR) techniques are powerful tools for this type of problems. IR techniques let find similar document based on different types of similarity measures with various similarity patterns including string matching, natural language processing and semantic analysis of text corpus.

¹Corresponding author, pedram@khu.ac.ir

Most simple plagiarism detection mechanisms only try to compare document texts using string or pattern matching techniques. These techniques are not using text's semantics in comparison, so they act very limited and cannot deal with synonym and polymorphisms (22), (8). Vector space models (VSM) have several advantages over these techniques. The models use vectors for both queries and documents, and words can be weighted. VSM models are able to use similarity measures such as cosine similarity, and they are often used in IR tasks such as clustering and classification (9).

However, these models are not able to deal with synonyms and polymorphisms. Synonym phrases have different dimensions in VSM and documents with similar words show a lower similarity measure respect to their real one. On the other hand document-term matrices are usually very large even for small data sets (10), (5). Singular Value Decomposition (SVD) algorithm can extract latent semantics from documents, thus SVD can be used to improve document similarity measures. While matrices resulting from decomposition are also very smaller than the original matrix (9), SVD is computationally expensive and this makes it insufficient for large data sets.

This paper uses natural language processing methods such as stop word removal and stemming to reduce vocabulary size of the corpus and using information retrieval methods such as clustering to improve its performance. In this paper, Latent Dirichlet Allocation (LDA) model is used to cluster documents and Natural Language Processing (NLP) techniques are applied to clean document texts.

2 Literature review

A plagiarism detection system, measures similarity of documents in a corpus and detects similar ones as suspicious documents (15). There are various measures and techniques used to measure this similarity and detect suspicious pairs of documents (16).

Character based techniques are the simplest and the most common techniques. These methods try to detect similarity between documents using string similarities (such as longest substring and frequent subset matching (21) or fingerprint based similarities (word n-gram and character n-gram hashing (20)). Though these techniques are simple to use and possess relatively good performance, they can't detect copied texts when author changes some word in copied text.

Syntax based methods are another category of plagiarism detection techniques. These methods are similar to character based techniques, but they also use natural language processing (NLP) algorithms generated metadata of words (such as part-of-speech tags (6) to find similar subsequences in different documents. These techniques are good at finding deeper level of similarities, but they are not able to use word synonymies in similarity measures.

Another category of techniques are semantic based techniques which use word similarity graphs such as WordNet (13) or embedding techniques such as LDA (2) in order to apply semantic similarity between words to better understand similarities between documents (4). Ceska (3) used Singular Value Decomposition (SVD) on corpora to build better similarity measure on word similarities to use these measures in document similarity detection.

There are several researches on applying plagiarism techniques on Persian language corpora. Mahmoudi used different types of similarity measures like Jackard distance and Clough-Stevenson measure on Persian documents (12). The paper also used a set of NLP techniques such as stop-word removal, stemming, lemmatization, part-of-speech tagging. Esteki (7) applied Support Vector Machine (SVM) on these measures to find a better way to combine these measures to find similar documents. Mahdavi and et al. in (11) used vector space model and cosine similarity to find similar documents. Rafieian tries fingerprinting based methods on Persian corpus (19).

3 Proposed method

Our proposed method consists of three major steps, which are shown in Figure 1. These three steps are described in this section.

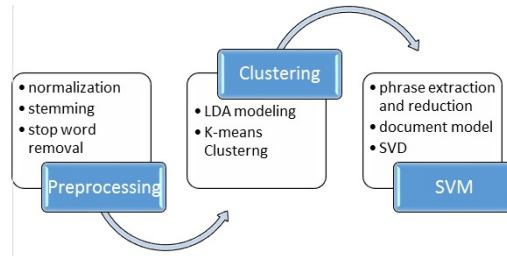


Figure 1: *Proposed method block diagram.*

3.1 Preprocessing

Preprocessing is a pivotal step in every natural language processing task. We first normalized our corpuses to remove excessive characters and unified different shape of characters. Then we removed stop words using a list of stop words (1).

Persian language processing has some difficulties with multiple part words (with spaces in between), different shapes for same word and different morphological rules for suffixes and prefixes (17). Thus we removed excessive characters and unified different shape of characters. Then verbs are stemmed using a reference collection and suffixes are then removed (1).

3.2 Clustering

We use a generative model named Latent Dirichlet Allocation to cluster documents based on topics discussed in each document. Generative models assume words of a document are generated from a probabilistic model (this method considers the documents as bag-of-words and does not use word ordering in its computations). The Goal of LDA is to find latent variables of these models with respect to words as visible variables (18). LDA assumes topics are distributions over words and documents are distributions over topics. Also LDA assumes documents have several topics (2).

We modeled our corpus using LDA, thus each document are represented as a probabilistic distribution over a set of topics. Documents are shown as a vector of topic weights. Similarity between two documents can be estimated by measuring the similarity between their distributions. Various methods can be used to cluster document distributions at this point. We used k-means clustering for this purpose.

3.3 Plagiarism detection

Ceska used SVD as a plagiarism detection method. First in each cluster all n-gram phrases are extracted. Phrases that only occur once and insignificant phrases are removed from each set. Then the document-phrase matrices are built (3)

$$M = U\Sigma V^T$$

V^T Is the document-factor matrix which is a set of Eigenvectors of documents. Eigenvectors of documents are then weighted with the corresponding eigenvalues:

$$B = \Sigma V^T$$

The matrix B is column normalized (3). Then, correlation matrix is built to measure similarity between documents:

$$C = B^T B$$

Each element of C shows similarity percentage between two documents (9). We applied this method to detect document pairs with higher similarity than a threshold t as plagiarized.

3.4 Experimental results

To evaluate our method, we use standard mechanism in IR systems. We define precision as number of relevant documents which are retrieved by the system on number of retrieved documents. And recall is the number of relevant documents which are retrieved on number of relevant documents.

$$Precision = \frac{TP}{(TP + FN)}$$
$$Recall = \frac{TP}{(TP + FP)}$$

Then we determine a practical threshold on F_1 measure:

$$F_1 = \frac{(2 \times Precision \times Recall)}{(Precision + Recall)}$$

Some experiments were conducted on two sets of data sets in English and Persian to show the effectiveness of the proposed method. For English data set we used 500 documents from Gothenburg project². We randomly selected some of these documents and made plagiarism cases using following rules and added these new documents to the corpora:

- 1 , 15% of documents have one source, 45% has two source, 25% has three sources and 15% has more than three sources.
- 2 , From each source some phrases and paragraphs are randomly picked and added to the text.
- 3 , Some phrases and paragraphs are randomly chosen and replaced by their synonyms.

For Persian data set we used 300 documents from the 11th Iranian conference on intelligent systems. We randomly selected some of these documents and made plagiarism cases using above rules. Then we added these new documents to the set.

The proposed method was tested under different thresholds and various n for n-grams. Figure 2 shows influence of thresholds on F_1 measure w.r.t. different n-grams on English data set.

²www.gutenberg.org

As shown in Figure 2, the best F_1 measure is 0.88 for 2-grams SVD clustering system w.r.t. 35 percent threshold. For 4-gram SVD we achieved 0.92 for F_1 w.r.t. 30 percent threshold and for 5-gram we have 0.93 w.r.t. 15 percent threshold.

Table 1 shows precision, recall and F_1 measure for English data set. As shown in this table best F_1 measure is achieved with 5-gram w.r.t. threshold 15 percent.

Table 1: F_1 measure for different n in n -grams on English data set

method	threshold	precision	recall	F_1
2-gram Cluster SVD	35	0.88	0.97	0.88
4-gram Cluster SVD	30	0.99	0.84	0.92
5-gram Cluster SVD	15	0.99	0.87	0.92

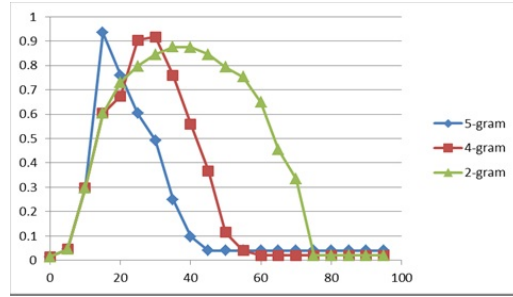


Figure 2: F_1 measure in English data set w.r.t. different n -grams

Figure 1 shows the influence of threshold on F_1 measure w.r.t. different n -grams on Persian data set. Best measure F_1 is 0.65 for 2-grams SVD clustering system w.r.t. 10 percent threshold. For 4-gram SVD we achieved 0.86 measure F_1 w.r.t. 10 percent threshold and for 5-gram we have 0.90 w.r.t. 4 percent threshold.

Table 2 shows precision, recall and F_1 measure for Persian data set. As shown in this table, best F_1 measure is achieved with 5-gram w.r.t. threshold 5 percent.

Table 2: F_1 measure for different n in n -grams on Persian data set

method	threshold	precision	recall	F_1
2-gram Cluster SVD	10	0.63	0.68	0.65
4-gram Cluster SVD	10	0.86	0.86	0.86
5-gram Cluster SVD	5	0.87	0.92	0.90

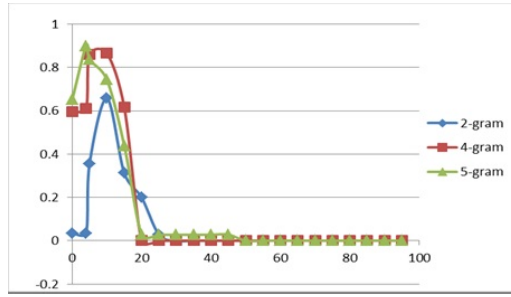


Figure 3: F_1 measure in Persian data set w.r.t. different n-grams

As shown in Figure 2 and Figure 1, with the lower n in n-grams, curves flatness is increased. Finding the right threshold is easier in these cases. However more unrelated documents are detected as plagiarized, so precision is reduced. On the other hand with higher n , curves are narrower, therefore the maximum value of F_1 measure is achieved on lower thresholds.

4 Conclusion

Experimental results show increasing n-gram length results in higher accuracy. This also increases the sensitivity to small changes in copied texts and method is not able to recognize the copied texts, although they only cause small variations in texts, which decreases recall.

Clustering the documents shows a significant increase in accuracy. It also helps to find plagiarized documents in large corpus faster and easier. One of the problems we faced in this research was lack of accurate Persian preprocessing tools. Our experiments show increasing preprocessing accuracy has significant effect on whole model's accuracy.

References

- [1] Asgari, E., & Chappelier, J.-C. (2013). Linguistic Resources and Topic Models for the Analysis of Persian Poems. *CLfL@ NAACL-HLT*.
- [2] Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent dirichlet allocation. *Journal of machine Learning research*, 3(Jan), 993-1022.
- [3] Ceska, Z. (2008). Plagiarism detection based on singular value decomposition. *Advances in natural language processing*, 108-119.
- [4] Ceska, Z., & Fox, C. (2011). The influence of text pre-processing on plagiarism detection. *Association for Computational Linguistics*.
- [5] Deerwester, S. e. (1990). Indexing by latent semantic analysis. *Journal of the American society for information science*.
- [6] Elhadi, M., & A. Al-Tobi. (2009). Duplicate detection in documents and webpages using improved longest common subsequence and documents syntactical structures. *Fourth International Conference on Computer Sciences and Convergence Information Technology*.
- [7] Esteki, F., & Esfahani, F. S. (2016). A Plagiarism Detection Approach Based on SVM for Persian Texts. In *FIRE (Working Notes)* (pp. 149-153).

- [8] Heintze, N. (1996). Scalable document fingerprinting. *USENIX workshop on electronic commerce*.
- [9] Hofmann, T. (2001). Unsupervised learning by probabilistic latent semantic analysis. In *Machine learning* (pp. 177-196).
- [10] Landauer, T., Foltz, P., & Laham, D. (1998). An introduction to latent semantic analysis. In *Discourse processes* (pp. 259-284).
- [11] Mahdavi, P., Siadati, Z., & Yaghmaee, F. (2014). Automatic external Persian plagiarism detection using vector space model. *Computer and Knowledge Engineering (ICCKE), 2014 4th International eConference on*.
- [12] Mahmoodi, M., & Varnamkhasti, M. M. (2014). Design a Persian Automated Plagiarism Detector (AMZPPD). *International Journal of Engineering Trends and Technology(IJETT)* , 8(8).
- [13] Miller, G., & Fellbaum, C. (1998). Wordnet: An electronic lexical database.
- [14] Nelson, R. S. (2007). In *Random House Dictionary of the English Language* (p. 65). Assoc. of College & Resrch Libraries.
- [15] Osman, A. H., Salim, N., & Abuobieda., A. (2012). Survey of text plagiarism detection. *Computer Engineering and Applications Journal (ComEngApp)* (1.1), 37-45.
- [16] Osman, A. H., Salima, N., & Binwahlan, M. S. (2012). An improved plagiarism detection scheme based on semantic role labeling. *Applied Soft Computing*, 1493-1502.
- [17] Dolamic, L., & Savoy, J. (2009, August). Persian language, is stemming efficient?. In *Database and Expert Systems Application, 2009. DEXA '09. 20th International Workshop on* (pp. 388-392). IEEE.
- [18] Steyvers, M., & Griffiths, T. (2007). Probabilistic topic models. *Handbook of latent semantic analysis*, 427(7), 424-440.
- [19] Rafieian, S. (2016). Plagiarism checker for Persian (PCP) texts using hash-based tree representative fingerprinting. *Journal of AI and Data Mining*, 4(2).
- [20] Sánchez-Vega, F., Villatoro-Tello, E., Montes-y-Gómez, M., Rosso, P., Stamatatos, E., & Villaseñor-Pineda, L. (2017). Paraphrase plagiarism identification with character-level features. *Pattern Analysis and Applications*, 1-13.
- [21] Shivakumar, N., & Garcia-Molina, H. (1995). SCAM: A copy detection mechanism for digital documents.
- [22] Stein, B., & Eissen, S. z. (2007). In *Fingerprint-based Similarity Search and its Applications*. Universität Weimar.

A capable neural network model for fuzzy quadratic optimization problems

Amin Mansoori¹ ¹, Sohrab Effati¹,

¹Department of Applied Mathematics, Ferdowsi University of Mashhad, Mashhad, Iran

Abstract

In this paper, a representation of a recurrent neural network to solve fuzzy quadratic programming problems (FQP) is given. The motivation of the paper is to design a new effective one-layer structure recurrent neural network model for solving the FQP. Here, we reformulate the FQP to a bi-objective problem. Furthermore, the bi-objective problem is reduced to a weighting problem and then the Karush-Kuhn-Tucker (KKT) optimality conditions of the problem are constructed. A novel recurrent neural network model to solve the FQP is developed. Finally, an illustrative example is presented.

Keywords: Fuzzy quadratic programming problem, recurrent neural network model, bi-objective problem, weighting problem, globally stable in the sense of Lyapunov.

1 Introduction

In this paper, we study the following FQP (quadratic programming problems with fuzzy parameters):

$$\begin{aligned} \min \quad & \tilde{f}(x) = \tilde{c}^T x + \frac{1}{2} x^T \tilde{H} x, \\ \text{s.t.} \quad & \tilde{A}x \leq \tilde{b}, \\ & x \geq 0, \end{aligned} \tag{1}$$

where $\tilde{H}$ is a fuzzy positive semi-definite and symmetric matrix (see Definition 2) of order $n \times n$, $\tilde{A}$, $\tilde{c}$, and $\tilde{b}$ are matrices of order $m \times n$, $n \times 1$, and $m \times 1$, respectively. Some real world problems and many interesting applications are classified as quadratic programming

¹Speaker: am.ma7676@yahoo.com; a-mansoori@um.ac.ir

problem such as inventory management (1), portfolio selection (2; 11), engineering design (9), etc. These real problems have uncertain and vague data that can be dealt with using fuzzy logic. Thus, the development of methods for solving quadratic programming problem under fuzzy environment emerges as a way of solving these kinds of problems.

Neural networks concepts were found in 1985 by Tank and Hopfield (5). There exist some neural network models for solving the quadratic optimization problem (7; 3; 4). Also, Mansoori et al. (8) proposed a recurrent neural network to solve the non-linear optimization problem with fuzzy parameters.

The motivation of this paper is to design a new effective one-layer structure neural network model for solving the FQP defined in (1). As mentioned above, there are several neural network models to solve crisp mathematical programming problems. This methodology is going to study the possibility of extending the existing neural network model to solve the FQP. Additionally, for showing the stability of the proposed neural network model, we give a Lyapunov function for dynamical system.

2 Preliminaries

In this section, we give some preliminaries of fuzzy set theory and optimization.

[(10)] A fuzzy number is a fuzzy set $\tilde{u} : \mathbb{R} \rightarrow [0, 1]$ with the following properties:

1. $\tilde{u}$ is upper semi-continuous function.
2. $\tilde{u}$ is normal, i. e., there exists an $x_0 \in \mathbb{R}$, with $\tilde{u}(x_0) = 1$.
3. $\tilde{u}$ is fuzzy convex, i. e., $\tilde{u}((1 - \lambda)x + \lambda y) \geq \min\{\tilde{u}(x), \tilde{u}(y)\}$ whenever $x, y \in \mathbb{R}$ and $\lambda \in [0, 1]$.
4. $cl(supp \tilde{u}) = cl\{x \in \mathbb{R} : \tilde{u}(x) > 0\}$ is a compact set.

The collection of all fuzzy numbers is denoted by E^1 . Also, the α -cut sets of a fuzzy number $\tilde{u} \in E^1$ is denoted by $[\underline{u}(\alpha), \bar{u}(\alpha)]$, that is defined by,

$$\tilde{u}[\alpha] = [u]^\alpha = [\underline{u}(\alpha), \bar{u}(\alpha)] = \begin{cases} \{x \in \mathbb{R} : u(x) \geq \alpha\}, & \text{if } 0 < \alpha \leq 1 \\ cl(supp \tilde{u}), & \text{if } \alpha = 0. \end{cases}$$

[(10)] A triangular fuzzy number is denoted by $\tilde{u} = (a, b, c)$, and its α -cuts is denoted by $\tilde{u}[\alpha] = [a + \alpha(b - a), c - \alpha(c - b)]$. [(10)]

The opposite of a fuzzy number $\tilde{u}$ to be the fuzzy number $-\tilde{u}$. If $\tilde{u}[\alpha] = [\underline{u}(\alpha), \bar{u}(\alpha)]$, then $-\tilde{u}[\alpha] = [-\bar{u}(\alpha), -\underline{u}(\alpha)]$. Also,

1. $(k\tilde{u})[\alpha] = [\min\{k\underline{u}(\alpha), k\bar{u}(\alpha)\}, \max\{k\underline{u}(\alpha), k\bar{u}(\alpha)\}], \quad \forall k \in \mathbb{R}^+.$
2. $(\tilde{u} + \tilde{v})[\alpha] = [\underline{u}(\alpha) + \underline{v}(\alpha), \bar{u}(\alpha) + \bar{v}(\alpha)].$
3. $(\tilde{u} - \tilde{v})[\alpha] = [\underline{u}(\alpha) - \bar{v}(\alpha), \bar{u}(\alpha) - \underline{v}(\alpha)].$
4. $(\tilde{u}\tilde{v})[\alpha] = [(\underline{uv})(\alpha), (\bar{uv})(\alpha)],$ where:

$$\begin{aligned} (\underline{uv})(\alpha) &= \min\{\bar{u}(\alpha)\underline{v}(\alpha), \bar{u}(\alpha)\bar{v}(\alpha), \underline{u}(\alpha)\underline{v}(\alpha), \underline{u}(\alpha)\bar{v}(\alpha)\}, \\ (\bar{uv})(\alpha) &= \max\{\bar{u}(\alpha)\underline{v}(\alpha), \bar{u}(\alpha)\bar{v}(\alpha), \underline{u}(\alpha)\underline{v}(\alpha), \underline{u}(\alpha)\bar{v}(\alpha)\}. \end{aligned} \tag{2}$$

[(10)] Let $f(x) = [\underline{f}(x), \overline{f}(x)]$ be an interval-valued function on $\Omega \subseteq \mathbb{R}^n$. We consider the following minimized problem with the interval-valued objective function:

$$\begin{aligned} \min \quad & f(x) = [\underline{f}(x), \overline{f}(x)], \\ \text{s.t.} \quad & x \in \Omega, \end{aligned} \tag{3}$$

where the feasible set Ω is assumed as a convex subset on $\mathbb{R}^n$. Consider the following bi-objective optimization problem:

$$\begin{aligned} \min \quad & (\underline{f}(x), \overline{f}(x)), \\ \text{s.t.} \quad & x \in \Omega. \end{aligned} \tag{4}$$

Also, consider the weighting problem of (4) defined by:

$$\begin{aligned} \min \quad & w_1 \underline{f}(x) + w_2 \overline{f}(x), \\ \text{s.t.} \quad & x \in \Omega, \end{aligned} \tag{5}$$

where w_1 and w_2 are positive real numbers and $w_1 + w_2 = 1$. It can be shown that, if x^* is an optimal solution of problem (5), then x^* is a Pareto optimal solution of problem (4) and also, x^* is an efficient solution of problem (3). [Positive semi-definite and symmetric fuzzy matrix] An $n \times n$ fuzzy matrix $\tilde{A} = (\{\underline{a}_{ij}(\alpha), \overline{a}_{ij}(\alpha), \alpha : \alpha \in [0, 1]\})_{n \times n}$ is said to be positive semi-definite and symmetric, if for any $\alpha \in [0, 1]$, $\underline{A}(\alpha) = (\underline{a}_{ij}(\alpha))_{n \times n}$ and $\overline{A}(\alpha) = (\overline{a}_{ij}(\alpha))_{n \times n}$ are $n \times n$ positive semi-definite and symmetric.

3 Fuzzy Quadratic Programming Problem

In this section, the FQP is presented. Let $\bar{x} \in \mathbb{R}_+^n$ is a local optimal solution to the problem (1), then $\bar{x}$ is a global optimal solution of (1). Also, if $\tilde{f}$ is strictly convex, then $\bar{x}$ is the unique global optimal solution of (1). Consider the FQP in (1). We can rewrite this problem as,

$$\begin{aligned} \min \quad & \tilde{f}(x) = \tilde{c}^T x + \frac{1}{2} x^T \tilde{H} x, \\ \text{s.t.} \quad & \tilde{A}'x \leq \tilde{b}', \end{aligned} \tag{6}$$

where $\tilde{A}' = \begin{bmatrix} \tilde{A} \\ -I \end{bmatrix}$, $\tilde{A}'(\alpha) = \begin{bmatrix} \underline{A}(\alpha) \\ \overline{A}(\alpha) \\ -I \end{bmatrix}$ and $\tilde{b}' = \begin{bmatrix} \tilde{b} \\ 0 \end{bmatrix}$, $\tilde{b}'(\alpha) = \begin{bmatrix} \underline{b}(\alpha) \\ \overline{b}(\alpha) \\ -I \end{bmatrix}$. Based on Definition 2, we can write (6) as an interval program:

$$\begin{aligned} \min \quad & [\underline{f}(x)(\alpha), \overline{f}(x)(\alpha)] \\ \text{s.t.} \quad & [\underline{A}'(\alpha), \overline{A}'(\alpha)]x \leq [\underline{b}'(\alpha), \overline{b}'(\alpha)], \\ & 0 \leq \alpha \leq 1, \end{aligned} \tag{7}$$

or it can be written as a bi-objective problem:

$$\begin{aligned}
\min \quad & (\underline{f}(x)(\alpha), \bar{f}(x)(\alpha)) \\
s.t. \quad & \underline{A}'(\alpha)x \leq \underline{b}'(\alpha), \\
& \overline{A}'(\alpha)x \leq \overline{b}'(\alpha), \\
& 0 \leq \alpha \leq 1,
\end{aligned} \tag{8}$$

where $\underline{f}(x)(\alpha) = \underline{c}(\alpha)^T x + \frac{1}{2} x^T \underline{H}(\alpha)x$ and $\bar{f}(x)(\alpha) = \bar{c}(\alpha)^T x + \frac{1}{2} x^T \overline{H}(\alpha)x$. By reformulating (8) into a weighting problem as in (5), we have:

$$\begin{aligned}
\min \quad & w_1 \underline{f}(x)(\alpha) + w_2 \bar{f}(x)(\alpha) \\
s.t. \quad & \underline{A}'(\alpha)x \leq \underline{b}'(\alpha), \\
& \overline{A}'(\alpha)x \leq \overline{b}'(\alpha), \\
& 0 \leq \alpha \leq 1,
\end{aligned} \tag{9}$$

where $w_1 + w_2 = 1$. The Lagrangian of (9) is derived as:

$$L(x, u_1, u_2) = w_1 \underline{f}(x)(\alpha) + w_2 \bar{f}(x)(\alpha) + u_1^T (\underline{A}'(\alpha)x - \underline{b}'(\alpha)) + u_2^T (\overline{A}'(\alpha)x - \overline{b}'(\alpha)), \quad u_1, u_2 \geq 0, \quad 0 \leq \alpha \leq 1.$$

The KKT optimality conditions for the problem (9) can be constructed as follow:

$$\text{KKT Conditions : } \begin{cases} u_1^*, u_2^* \geq 0, \quad \underline{A}'(\alpha)x^* - \underline{b}'(\alpha) \leq 0, \quad \overline{A}'(\alpha)x^* - \overline{b}'(\alpha) \leq 0, \\ u_1^{*T} (\underline{A}'(\alpha)x^* - \underline{b}'(\alpha)) = 0, \quad u_2^{*T} (\overline{A}'(\alpha)x^* - \overline{b}'(\alpha)) = 0, \\ = 0. \end{cases} \tag{10}$$

x^* is an optimal solution of the problem (1), if and only if, x^* satisfies the KKT conditions (10).

4 A Recurrent Neural Network Model

In this section, our aim is to construct a continuous-time dynamical system that will settle down to the KKT point of FQP (1).

Now, we may describe the recurrent neural network model corresponding to (1) can be proposed as the following non-linear dynamical system:

$$\frac{dx_\alpha}{dt} = -\nabla_x L(x, u_1, u_2) = -\left(w_1(\underline{c}(\alpha) + \underline{H}(\alpha)x) + w_2(\bar{c}(\alpha) + \overline{H}(\alpha)x) + u_1^{*T} \underline{A}'(\alpha) + u_2^{*T} \overline{A}'(\alpha)\right), \tag{11}$$

$$\frac{du_{1\alpha}}{dt} = \nabla_{u_1} L(x, u_1, u_2) = \text{diag}(u_{11}, \dots, u_{1m})(\underline{A}'(\alpha)x - \underline{b}'(\alpha)), \tag{12}$$

$$\frac{du_{2\alpha}}{dt} = \nabla_{u_2} L(x, u_1, u_2) = \text{diag}(u_{21}, \dots, u_{2m})(\overline{A}'(\alpha)x^* - \overline{b}'(\alpha)), \tag{13}$$

with an initial point (x_0, u_{10}, u_{20}) and $u_{1k}(t_0) \neq 0, u_{2k}(t_0) \neq 0, (k = 1, 2, \dots, m)$. Note that, $u_1 = (u_{11}, \dots, u_{1m}), u_2 = (u_{21}, \dots, u_{2m})$ are the Lagrange multipliers vectors, and also $\frac{d\alpha}{dt}$ means for any fixed α we have a particular dynamical system which refers to α . Now, let $y = (x, u_1, u_2) \in \mathbb{R}^{n+2m}$ be the optimal point of (1) and consider

$$\Phi(x, u_1, u_2) = \begin{bmatrix} -\left(w_1(\underline{c}(\alpha) + \underline{H}(\alpha)x) + w_2(\bar{c}(\alpha) + \overline{H}(\alpha)x) + u_1^{*T} \underline{A}'(\alpha) + u_2^{*T} \overline{A}'(\alpha)\right) \\ \text{diag}(u_{11}, \dots, u_{1m})(\underline{A}'(\alpha)x - \underline{b}'(\alpha)) \\ \text{diag}(u_{21}, \dots, u_{2m})(\overline{A}'(\alpha)x^* - \overline{b}'(\alpha)) \end{bmatrix}.$$

Therefore, the recurrent neural network model (11)-(13) can be summarized as the following:

$$\frac{dy}{dt} = \gamma \Phi(y), \quad y(t_0) = y_0, \quad u_1(t_0) \neq 0, \quad u_2(t_0) \neq 0, \quad (14)$$

where $\gamma > 0$ is the convergence rate.

5 Stability and Convergence Analysis

In this section, we show that the recurrent neural network model whose dynamics is described by the differential equation (14) is stable.

First, some definitions from dynamical system are presented (6).

[Stability in the sense of Lyapunov] Equilibrium point x^* is said to be stable in the sense of Lyapunov, if for any $y_0 = y(t_0)$ and any $\varepsilon > 0$, there exists a $\delta > 0$, such that:

$$\|y(t) - y^*\| < \varepsilon, \quad \forall t \geq t_0, \quad \|y(t_0) - y^*\| < \delta.$$

[Asymptotic stability] An isolated equilibrium point x^* is said to be asymptotically stable, if in addition to being Lyapunov stable, it has the property that $y(t) \rightarrow y^*$ as $t \rightarrow \infty$ for all $\|y(t_0) - y^*\| < \delta$.

Assume that, $y = (x^*, u_1^*, u_2^*)$ is the equilibrium of the recurrent neural network model (14). Then x^* is a KKT point of the problem (1). The equilibrium point of the proposed recurrent neural network model (14) is unique. There exists a unique solution $y(t)$ for the recurrent neural network (14). The proposed recurrent neural network model (14) is globally stable in the sense of Lyapunov and is globally convergent to the optimal solution of the problem (1).

6 Numerical Example

In this section, we present an illustrative example in order to demonstrate the applicability and the performance of the recurrent neural network model. Consider the following FQP:

$$\begin{aligned} \min \quad & \tilde{f} = (-6, -5, -4)x_1 + (1, 1.5, 2)x_2 + \frac{1}{2}[(4, 6, 8)x_1^2 + (-6, -4, -2)x_1x_2 + (2, 4, 6)x_2^2] \\ \text{s.t.} \quad & x_1 + (0.5, 1, 1.5)x_2 \leq (1, 2, 3) \\ & (1, 2, 3)x_1 + (-2, -1, -0.5)x_2 \leq (3, 4, 5) \\ & x_1, x_2 \geq 0, \end{aligned}$$

According to the α -cuts of fuzzy coefficients, we have:

$$\begin{aligned} \underline{H}(\alpha) &= \begin{bmatrix} 4 + 2\alpha & -3 + \alpha \\ -3 + \alpha & 2 + 2\alpha \end{bmatrix}, & \underline{c}(\alpha) &= \begin{bmatrix} -6 + \alpha \\ 1 + 0.5\alpha \end{bmatrix}, & \underline{A}(\alpha) &= \begin{bmatrix} 1 & 0.5 + 0.5\alpha \\ 1 + \alpha & -2 + \alpha \end{bmatrix}, & \underline{b}(\alpha) &= \begin{bmatrix} 1 + \alpha \\ 2 + 2\alpha \end{bmatrix}, \\ \overline{H}(\alpha) &= \begin{bmatrix} 8 - 2\alpha & -1 - \alpha \\ -1 - \alpha & 6 - 2\alpha \end{bmatrix}, & \overline{c}(\alpha) &= \begin{bmatrix} -4 - \alpha \\ 2 - 0.5\alpha \end{bmatrix}, & \overline{A}(\alpha) &= \begin{bmatrix} 1 & 1.5 - 0.5\alpha \\ 3 - \alpha & 0.5 - 0.5\alpha \end{bmatrix}, & \overline{b}(\alpha) &= \begin{bmatrix} 3 - \alpha \\ 6 - 2\alpha \end{bmatrix}. \end{aligned}$$

The primal and the dual solutions for various values of α , $w_1 = \frac{1}{3}$, and $w_2 = \frac{2}{3}$ are shown in Table 1. One can see all simulation results demonstrate that the neural network (14) is globally asymptotically stable. We solve this problem by letting initial point $(-3, -1.5, -0.5, 1, 2, 3)$. The α -cuts of f represent the possibility that the objective values will appear in the associated range. For instance, at $\alpha = 0.4$, the value of $f = -1.8000$ occurs at $x_1 = 0.7500$ and $x_2 = 0.0000$ and at $\alpha = 0.6$, the value of $f = -1.8897$ occurs at $x_1 = 0.7766$ and $x_2 = 0.0000$. Also, Figure 1 displays the transient behaviour of state trajectories based on (14).

Table 1: The solution of neural network (14) for different values of α , $w_1 = \frac{1}{3}$, and $w_2 = \frac{2}{3}$ in Example 6.

α	f	x_1	x_2	u_1	u_2	u_3	u_4	u_5	u_6	CPU-Time (s)
0.0	-1.6333	0.7000	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.5000	0.033967
0.2	-1.7146	0.7245	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.3776	0.030143
0.4	-1.8000	0.7500	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.2500	0.032257
0.5	-1.8443	0.7632	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.1842	0.032965
0.6	-1.8897	0.7766	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.1170	0.035451
0.8	-1.9841	0.8063	0.0062	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.035251
1.0	-2.0875	0.8500	0.0500	0.0000	0.0000	0.0000	0.0000	0.0000	0.0000	0.034296

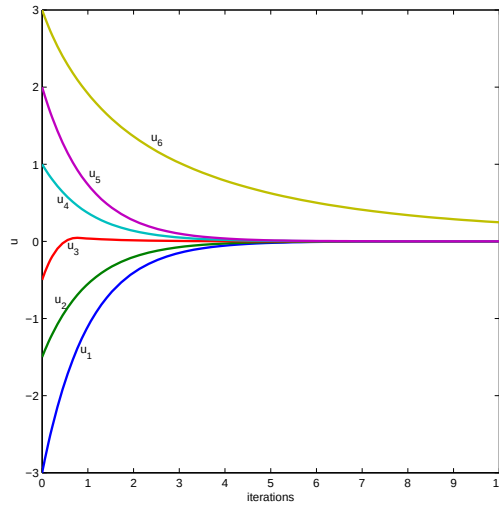


Figure 1: Transient behaviour of the neural network (14) for $\alpha = 0.6$, $w_1 = \frac{1}{3}$, and $w_2 = \frac{2}{3}$ in Example 6.

Conclusion

In this paper, an efficient recurrent neural network model for solving FQP was proposed. We solved the proposed recurrent neural network model by an ODE. Here, we reformulated the KKT conditions of the FQP into an efficient recurrent neural network model to solve the FQP. Additionally, by using these transformations, we found both the primal and the dual solutions of the original problem.

References

- [1] L. L. Abdel-Malek and N. Areeractch, A quadratic programming approach to the multi-product newsvendor problem with side constraints, *European Journal of Operational Research* **176** (8) (2007), 55-61.
- [2] E. Ammar and H. A. Khalifa, Fuzzy portfolio optimization a quadratic programming approach, *Chaos, Solitons and Fractals* **18** (2003), 1045-1054.
- [3] S. Effati, A. Mansoori, and M. Eshaghezhad, An efficient projection neural network for solving bilinear programming problems, *Neurocomputing* **168** (2015), 1188-1197.

- [4] M. Eshaghnezhad, S. Effati, and A. Mansoori, A Neurodynamic Model to Solve Nonlinear Pseudo-Monotone Projection Equation and Its Applications, *IEEE Transactions on Cybernetics* **47** (10) (2016), 3050-3062.
- [5] J. J. Hopfield and D. W. Tank, Neural computation of decisions in optimization problems, *Biological Cybernetics* **52** (1985), 141-152.
- [6] H. K. Khalil, *Nonlinear Systems*, Prentice-Hall, Michigan, NJ, 1996.
- [7] A. Mansoori, S. Effati, and M. Eshaghnezhad, A neural network to solve quadratic programming problems with fuzzy parameters, *Fuzzy Optimization and Decision Making* **17** (1) (2018), 75-101.
- [8] A. Mansoori, S. Effati, and M. Eshaghnezhad, An efficient recurrent neural network model for solving fuzzy non-linear programming problems, *Applied Intelligence* **46** (2016), 308-327.
- [9] J. A. M. Petersen, M. Bodson, Constrained quadratic programming techniques for control allocation, *IEEE Transactions on Control Systems Technology* **14** (9) (2006), 1-8.
- [10] H. -C. Wu, Evaluate fuzzy optimization problems based on biobjective programming problems, *Computers and Mathematics with Applications* **47** (2004), 893-902.
- [11] X. -L. Wu and Y. -K. Liu, Optimizing fuzzy portfolio selection problems by parametric quadratic programming, *Fuzzy Optimization and Decision Making* **11** (4) (2012), 411-449.

Fuzzy Ridge Linear Regression Analysis for Fuzzy Input–Output Data

Mohammad Reza Rabiei^{1 1} , Mohammad Arashi^{2 1} ,

Vahid Goodarzi^{3 1} , Fatemeh Piyadeh Kouhsar^{4 2} ,

¹Department of Statistics, Faculty of Mathematical Sciences, Shahrood University of Technology, Shahrood, 3619995161 Iran

² Department of Statistics, School of Mathematical Sciences, Ferdowsi University of Mashhad, Mashhad, 91775 Iran

Abstract

In this paper, we propose a new weighted fuzzy ridge regression model for a given set of fuzzy inputs and their corresponding triangular fuzzy outputs. Coefficients are also considered to be fuzzy. First, some approximations for multiplication of two triangular fuzzy numbers are introduced. Then, respective definitions of fuzzy weighted arithmetic (FWA) are given. Followingly, a method for fuzzy least-squares ridge regression based on FWA is developed. To obtain the ridge tuning parameter an objective function is optimized. Lastly, for the illustration propose, a real example is analyzed, which indicates the superior performance of the proposed method.

Keywords: Multicollinearity, Ridge regression, Fuzzy regression, Weighted fuzzy arithmetic.

1 Introduction

Fuzzy regression analysis is a powerful tool for analyzing phenomena in which one variable namely output (response or dependent) variable depends on one or more variables namely input (explanatory or independent) variables, in a fuzzy environment. A fuzzy linear regression model was first proposed by (14) to analyze fuzzy data with a vague relationship between the dependent variable and

¹Mohammad Reza Rabiei: rabiei1354@yahoo.com

²m_arashi_stat@yahoo.com

³v1992gi@gmail.com

⁴f.kouhsar@gmail.com

independent variables. However, when multicollinearity exists among input variables, the fuzzy linear regression fails to perform well. The method of ridge regression, proposed by (9) is one of the most widely used tools to the problem of multicollinearity. In this paper, we will dealing with a ridge regression methodology where the data include the fuzzy inputs, output, and coefficients.

In this paper, we focus on ridge regression modeling in a fuzzy environment using the concept of a fuzzy weighted arithmetic. Making use of the least-squares method (2) and (1) proposed a new hybrid least-squares regression method for normalized triangular fuzzy outputs. We develop the later approach to the case where the inputs and coefficients are also fuzzy numbers. The rest of this paper is organized as follows. Section 2, illustrates preliminaries and ridge regression procedures for fuzzy multivariable linear models. In section 3, the weighted fuzzy arithmetic is introduced. In section 4, we propose a new fuzzy ridge regression approach to linear models. This approach is evaluated using the Portland cement data in section 5. We conclude the paper in section 6.

2 Preliminaries

In this section, we review some preliminary definitions and a well-known result of the triangular fuzzy numbers which is needed for the construction of fuzzy ridge regression. Let $\mathbb{R}$ be the real number set.

1.2. (11) A fuzzy set $\tilde{A}$ on $\mathbb{R}^1$ is called a fuzzy number iff $\tilde{A}$ is convex and there exists exactly one point, say $M \in \mathbb{R}^1$, with $m_{\tilde{A}}(M) = 1$ that $m_{\tilde{A}}(\cdot)$ is a membership function.

2.2. (13) Let $\tilde{A}$ and $\tilde{B}$ be two fuzzy sets. $\tilde{A}$ is said to be a subset of $\tilde{B}$; denoted by $\tilde{A} \subseteq \tilde{B}$; iff its membership function is less than or equal to that of $\tilde{B}$ everywhere on X

$$\tilde{A} \subseteq \tilde{B} \Leftrightarrow \tilde{A}(x) \leq \tilde{B}(x), \quad \forall x \in X. \quad (1)$$

3.2. (12) Suppose $\tilde{A}$, $\tilde{B}$, and $\tilde{C}$ are fuzzy numbers. Then

$$(\tilde{A} \oplus \tilde{B}) \otimes \tilde{C} \subseteq \tilde{A} \otimes \tilde{C} \oplus \tilde{B} \otimes \tilde{C}. \quad (2)$$

The equality (2) holds if $\tilde{A}$, $\tilde{B}$, $\tilde{C} \succ 0$.

4.2. (19) A triangular fuzzy number (TFN), say $\tilde{A}$, is denoted by $\tilde{A} = (a, \alpha, \beta)_T$ where a, α and β are the center, left, and right spread of $\tilde{A}$, respectively a is a real number and $\alpha, \beta > 0$. The membership function of $\tilde{A}$ is as follows

$$\tilde{A}(x) = \begin{cases} \frac{x - (a - \alpha)}{\alpha}, & a - \alpha \leq x \leq a, \\ \frac{a + \beta - x}{\beta}, & a \leq x \leq a + \beta, \\ 0, & \text{otherwise.} \end{cases} \quad (3)$$

Note that the lower and upper limits of the fuzzy number $\tilde{A}$ are $(a - \alpha)$ and $(a + \beta)$, respectively. We denote the set of all triangular fuzzy numbers defined on the real line $\mathbb{R}$ by $T(\mathbb{R})$.

5.2. The closed intervals $[A^L(h), A^U(h)] = [a - (1 - h)\alpha, a + (1 - h)\beta]$ and $[a - \alpha, a + \beta]$ are called h -level set $h \in (0, 1]$ and support of $\tilde{A}$; respectively. If $\alpha = \beta$; then $\tilde{A}$ is a symmetric TFN and is denoted by $\tilde{A} = (a, \alpha)_T$. A real number can be considered as a degenerated TFN whose spreads are zero.

6.2. The following formulas for addition of two TFNs and multiplication of a TFN by a scalar are drawn from extension principle of (18). If $\tilde{A}_1 = (a_1, \alpha_1, \beta_1)_T$ and $\tilde{A}_2 = (a_2, \alpha_2, \beta_2)_T$ be in $T(\mathbb{R})$ and $r \in \mathbb{R}$, then

$$\tilde{A}_1 \oplus \tilde{A}_2 = (a_1 + a_2, \alpha_1 + \alpha_2, \beta_1 + \beta_2)_T, \quad (4)$$

$$r\tilde{A}_1 = \begin{cases} (ra_1, r\alpha_1, r\beta_1)_T, & r \geq 0, \\ (ra_1, -r\beta_1, -r\alpha_1)_T, & r < 0. \end{cases} \quad (5)$$

7.2. (3): The multiplication of two TFNs based on extension principle is not necessarily a TFN. The following approximation is used for the product of two TFNs $\tilde{A}_1 = (a_1, \alpha_1, \beta_1)_T$ and $\tilde{A}_2 = (a_2, \alpha_2, \beta_2)_T$

$$\tilde{A}_1 \otimes \tilde{A}_2 \simeq \begin{cases} (a_1a_2, a_1\alpha_2 + a_2\alpha_1, a_1\beta_2 + a_2\beta_1)_T, & \tilde{A}_1 \succ 0, \tilde{A}_2 \succ 0, \\ (a_1a_2, a_2\alpha_1 - a_1\beta_2, a_2\beta_1 - a_1\alpha_2)_T, & \tilde{A}_1 \prec 0, \tilde{A}_2 \succ 0, \\ (a_1a_2, -a_2\beta_1 - a_1\beta_2, -a_2\alpha_1 - a_1\alpha_2)_T, & \tilde{A}_1 \prec 0, \tilde{A}_2 \prec 0, \end{cases} \quad (6)$$

where $\tilde{A}_i \succ 0$ ($\prec 0$) means $a_i - \alpha_i \geq 0$ ($a_i + \beta_i \leq 0$), $i = 1, 2$.

8.2. (4) Suppose we are given with the training data, $(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_n, y_n)$, where $\mathbf{x}_i$ are vectors in $\mathbb{R}^n$ and $y_i \in \mathbb{R}, i = 1, \dots, n$. The ridge regression procedure is a slight modification on the least-squares method and replaces the objective function $L(\mathbf{A})$ by

$$L(\mathbf{A}) = \sum_{i=1}^n (y_i - \mathbf{A} \cdot \mathbf{x}_i)^2 + k \|\mathbf{A}\|^2, \quad (7)$$

where the dot "." denotes the dot product in $\mathbb{R}^n$. The parameter k controls the smoothness and degree of fit. If there happens multicollinearity in the input variables, numerical stability problems occur and require the use of ridge regression.

3 Fuzzy Weighted Arithmetic

The fuzzy weighted arithmetic (FWA) defines the arithmetic operations between two fuzzy numbers as operating two corresponding values in each fuzzy set at the same membership level, integrating each level operation weighted by the membership level for the entire fuzzy sets, and dividing the weighted integration by the total integral of the membership function. The weighted fuzzy-arithmetic uses the concept of defuzzification to convert the operation of fuzzy sets into a crisp real number. The resulting crisp number can be interpreted as a mean value of a fuzzy-arithmetic operation. According to the definition of FWA, formulas for fuzzy-arithmetic operations are derived in this section. In order to perform the integration of the membership functions, fuzzy numbers are expressed as normalized asymmetrical triangular membership functions. If fuzzy numbers are not normalized, we normalize them through scaling with respect to the maximum membership value. Therefore, it can be proved that the total integral of a normalized triangular membership function is a unity.

1.3. (2) For any fuzzy number $\tilde{A} = (a_A, \alpha_A, \beta_A)_T$ in $T(\mathbb{R})$ and $h \in [0, 1]$ h -level set of $\tilde{A}$ is defined to be the closed interval

$$[\tilde{A}^L(h), \tilde{A}^U(h)] = [a_A - (1 - h)\alpha_A, a_A + (1 - h)\beta_A].$$

The weighted sum of any two fuzzy numbers $\tilde{A} = (a_A, \alpha_A, \beta_A)_T$ and $\tilde{B} = (a_B, \alpha_B, \beta_B)_T$ in $T(\mathbb{R})$ is computed as

$$\tilde{A} \oplus \tilde{B} = \frac{\int_0^1 (\tilde{A}^L(h) + \tilde{B}^L(h))h dh + \int_0^1 (\tilde{A}^U(h) + \tilde{B}^U(h))h dh}{\int h dh}, \quad (8)$$

where the denominator is taken to be $\int h dh = 2 \int_0^1 h dh = 1$. Thus, we have

$$\tilde{A} \oplus \tilde{B} = a_A + a_B + \frac{1}{6}((\beta_A + \beta_B) - (\alpha_A + \alpha_B)). \quad (9)$$

If both $\tilde{A}, \tilde{B} \in T(\mathbb{R})$ are symmetric numbers then $\alpha_A = \beta_A$ and $\alpha_B = \beta_B$ will be true, and in this case we get $\tilde{A} \oplus \tilde{B} = a_A + a_B$. Similarly, the weighted fuzzy subtraction can be obtained as $\tilde{A} \ominus \tilde{B} = (a_A - a_B) + \frac{1}{6}((\beta_A - \beta_B) - (\alpha_A - \alpha_B))$ and when both $\tilde{A}, \tilde{B}$ are symmetric numbers, it becomes $a_A + a_B$.

2.3. (17) The norm of a fuzzy number $\tilde{A} = (a_A, \alpha_A, \beta_A)_T$ in $T(\mathbb{R})$, denoted by $||\cdot||$, is computed by

$$||(a_A, \alpha_A, \beta_A)_T||^2 = a_A^2 + \frac{1}{3}a_A(\beta_A - \alpha_A) + \frac{1}{12}(\alpha_A^2 + \beta_A^2). \quad (10)$$

Since the weighted distance between the observed and the computed values are required for calculating the residual error in least-squares approach, we define the square of the weighted distance $d(\cdot, \cdot)$ between two given fuzzy numbers $\tilde{A} = (a_A, \alpha_A, \beta_A)_T$ and $\tilde{B} = (a_B, \alpha_B, \beta_B)_T$ in $T(\mathbb{R})$ using the norm of a fuzzy number defined by (10) as:

$$d^2(\tilde{A}, \tilde{B}) = (a_A - a_B)^2 + \frac{1}{3}(a_A - a_B)[(\beta_A - \beta_B) - (\alpha_A - \alpha_B)] + \frac{1}{12}[(\alpha_A - \alpha_B)^2 + (\beta_A - \beta_B)^2]. \quad (11)$$

4 The Proposed Method

Now, consider a set of TFN data $\{(\tilde{x}_{i1}, \tilde{x}_{i2}, \dots, \tilde{x}_{ip}, \tilde{y}_i) | i = 1, \dots, n\}$, in which $\tilde{x}_{ij} (i = 1, \dots, n, j = 1, \dots, p)$ is the value of j th independent variable and $\tilde{y}_i (i = 1, \dots, n)$ is the corresponding value of dependent variable $\tilde{\mathbf{y}}$, in the i th case. The purpose of fuzzy linear regression (FLR) is to fit a fuzzy linear model to the given TFN data. This model can be considered as follows

$$\tilde{\mathbf{Y}} = \tilde{A}_0 \oplus \tilde{A}_1 \otimes \tilde{\mathbf{x}}_1 \oplus \tilde{A}_2 \otimes \tilde{\mathbf{x}}_2 \oplus \dots \oplus \tilde{A}_p \otimes \tilde{\mathbf{x}}_p. \quad (12)$$

In the model (12), $\tilde{A}_0, \tilde{A}_1, \dots, \tilde{A}_p$ are FLR coefficients (parameters) which are fuzzy numbers. In particular, these parameters must be estimated so that the estimated responses $\tilde{Y}_i$,

$$\tilde{Y}_i = \tilde{A}_0 \oplus \tilde{A}_1 \otimes \tilde{x}_{i1} \oplus \tilde{A}_2 \otimes \tilde{x}_{i2} \oplus \dots \oplus \tilde{A}_p \otimes \tilde{x}_{ip}, \quad i = 1, \dots, n, \quad (13)$$

have the best fitness to the observed responses $\tilde{Y}_i (i = 1, \dots, n)$, according to a criterion of goodness of fit. Now, suppose that both inputs and outputs are non-symmetric TFN in the FLR model. Also, suppose that given inputs, $\tilde{x}_{ji} \ i = 1, \dots, n, j = 1, \dots, p$, are positive non-symmetric TFNs, by a simple translation of all data, if necessary. On the other hand, suppose that the observed responses, $\tilde{y}_i \in T(\mathbb{R}) (i = 1, \dots, n, j = 1, \dots, p)$, are non-symmetric TFNs. To estimate the regression coefficients, it is necessary to evaluate the product of two TFNs $\tilde{A}_j$ and $\tilde{x}_j$ in the model (12). It is clear that the multiplication $\tilde{A}_1 \otimes \tilde{A}_2$ in (6) depends on the sign of TFNs. Therefore, one must formulate different models for different states of the sign of regression coefficients. Another problem with using the multiplication (6) is that it is proposed only for two positive and/or negative TFNs. So, it cannot be used for TFNs whose supports contain both positive and negative real numbers. To handle this issue, we calculate the multiplication of two TFNs by some (in fact by infinite) TFNs. Then, the best approximation is chosen among them, to minimize a suitable function.

Note that a real number has infinite representations in the form of $a = a' - a''$, where a' and a'' are non-negative real numbers. Assume $\tilde{A} = (a, \alpha, \beta)_T$ and $\tilde{x} = (x, \alpha_x, \beta_x)_T$ are two TFNs where $\tilde{A}$ is unrestricted and $\tilde{x} \succ 0$. Set $\tilde{A} = \tilde{A}' \oplus \tilde{A}''$ in which $\tilde{A}' = (a', \alpha', \beta')_T$

and $\tilde{A}'' = (-a'', \alpha'', \beta'')_T$ where $a = a' - a'', \alpha = \alpha' + \alpha'', \beta = \beta' + \beta''$ and $a' \geq \alpha', a'' \geq \beta''$. By using Proposition 3.2, multiplication (6), and addition (4) we have

$$\begin{aligned}\tilde{A} \otimes \tilde{x} &= ((a', \alpha', \beta')_T \oplus (-a'', \alpha'', \beta'')_T) \otimes \tilde{x} \\ &\subseteq (a', \alpha', \beta')_T \otimes (x, \alpha_x, \beta_x)_T \oplus (-a'', \alpha'', \beta'')_T \otimes (x, \alpha_x, \beta_x)_T \\ &\simeq (a'x, a'\alpha_x + \alpha'x, a'\beta_x + \beta'x)_T \oplus (-a''x, a''\beta_x + \alpha''x, a''\alpha_x + \beta''x)_T.\end{aligned}\quad (14)$$

Therefore, $\tilde{A} \otimes \tilde{x}$ can be approximated as the RHS of (14) as

$$\tilde{A} \otimes \tilde{x} \simeq [(a' - a'')x, (a'\alpha_x + \alpha'x + a''\beta_x + \alpha''x), (a'\beta_x + \beta'x + a''\alpha_x + \beta''x)]. \quad (15)$$

Let $\tilde{A}_j = \tilde{A}'_j \oplus \tilde{A}''_j$ for $j = 1, \dots, p$, where $\tilde{A}'_j = (a'_j, \alpha'_j, \beta'_j)_T$, $\tilde{A}''_j = (-a''_j, \alpha''_j, \beta''_j)_T$, and $\tilde{A}'_j \succ 0, \tilde{A}''_j \prec 0$. For each choice of $\tilde{A}'_j$ and $\tilde{A}''_j$, by using approximation (15) and addition (4), we have the following approximation for i th estimated response $\tilde{Y}_i, i = 1, \dots, n$

$$\begin{aligned}\tilde{Y}_i &= \tilde{A}_0 \oplus \tilde{A}_1 \otimes \tilde{x}_{i1} \oplus \dots \oplus \tilde{A}_p \otimes \tilde{x}_{ip} \\ &\simeq [a_0 + \sum_{j=1}^p (a'_j - a''_j)x_{ij}, \alpha_0 + \sum_{j=1}^p (a'_j\alpha_{x_{ij}} + \alpha'_jx_{ij} + a''_j\beta_{x_{ij}} + \alpha''_jx_{ij}) \\ &\quad , \beta_0 + \sum_{j=1}^p (a'_j\beta_{x_{ij}} + \beta'_jx_{ij} + a''_j\alpha_{x_{ij}} + \beta''_jx_{ij})].\end{aligned}\quad (16)$$

Let us denote the RHS of approximation (16) by $\tilde{Y}_i(\tilde{A}', \tilde{A}'')$, so that it is a function of $\tilde{A}'$ and $\tilde{A}''$, which are m -dimensional vectors with j th elements $\tilde{A}'_j$ and $\tilde{A}''_j$, respectively. As we know, Y_i is not necessarily a TFN. However, for each choice of $\tilde{A}'$ and $\tilde{A}''$, $\tilde{Y}_i(\tilde{A}', \tilde{A}'')$ is a triangular approximation of $\tilde{Y}_i$ which is independent of the sign of regression coefficients. We consider it as an approximation for the i th response $\tilde{Y}_i$. Therefore, there are many approximations for $\tilde{Y}_i$, among which, we try to choose the best approximation. To this end, in a similar fashion as in (8), we attempt to make the membership function of each approximated response $\tilde{Y}_i(\tilde{A}', \tilde{A}'')$ as close as possible to the membership function of corresponding observed response $\tilde{Y}_i$. Hence, we solve a linear programming problem which finds the best choices of $\tilde{A}'$ and $\tilde{A}''$ for the coefficients of the model (12) by minimizing the total distance between the observed and approximated responses. Suppose that the observed responses are non-symmetric TFNs $\tilde{y}_i = (y_i, \alpha_{y_i}, \beta_{y_i})_T, i = 1, \dots, n$. The ridge regression problem (7) can be formulated as an optimization problem defined by

$$\text{Minimize } Q(A_0, A_1, \dots, A_p) = k(\|\tilde{A}_0\|^2 + \|\tilde{A}\|^2) + \sum_{i=1}^n d^2(\tilde{Y}_i(\tilde{A}', \tilde{A}''), \tilde{y}_i), \quad (17)$$

$$\begin{aligned}\text{s.t. } & a'_j \geq \alpha'_j, a''_j \geq \beta''_j \quad j = 1, \dots, p, \\ & a_0 \in \mathbb{R}, \alpha_0, \beta_0 \geq 0, \text{ and } a'_j, a''_j, \alpha'_j, \alpha''_j, \beta'_j, \beta''_j \geq 0, \quad j = 1, \dots, p,\end{aligned}\quad (18)$$

where

$$\|A\|^2 = \sum_{j=1}^p \|A_j\|^2 = \sum_{j=1}^p \|(a'_j - a''_j, \alpha'_j + \alpha''_j, \beta'_j + \beta''_j)\|^2. \quad (19)$$

On the other hand, for $i = 1, \dots, n$, we have

$$\begin{aligned}
d^2(\tilde{Y}_i(\tilde{A}', \tilde{A}''), \tilde{y}_i) &= [a_0 + \sum_{j=1}^p (a'_j - a''_j)x_{ij} - y_i]^2 + \frac{1}{3}(a_0 + \sum_{j=1}^p (a'_j - a''_j)x_{ij} - y_i) \\
&\quad \times [(\beta_0 + \sum_{j=1}^p (a'_j\beta_{x_{ij}} + \beta'_jx_{ij} + a''_j\alpha_{x_{ij}} + \beta''_jx_{ij}) - \beta_{y_i})] \\
&\quad - (\alpha_0 + \sum_{j=1}^p (a'_j\alpha_{x_{ij}} + \alpha'_jx_{ij} + a''_j\beta_{x_{ij}} + \alpha''_jx_{ij}) - \alpha_{y_i}) \\
&\quad + \frac{1}{12}[(\beta_0 + \sum_{j=1}^p (a'_j\beta_{x_{ij}} + \beta'_jx_{ij} + a''_j\alpha_{x_{ij}} + \beta''_jx_{ij}) - \beta_{y_i})^2 \\
&\quad + (\alpha_0 + \sum_{j=1}^p (a'_j\alpha_{x_{ij}} + \alpha'_jx_{ij} + a''_j\beta_{x_{ij}} + \alpha''_jx_{ij}) - \alpha_{y_i})^2].
\end{aligned} \tag{20}$$

Hence, we arrive at the following convex optimization problem for model (17)

$$\begin{aligned}
\text{Minimize } Q(A_0, A_1, \dots, A_p) &= k \left([a_0^2 + \frac{1}{3}a_0(\beta_0 - \alpha_0) + \frac{1}{12}(\alpha_0^2 + \beta_0^2)] \right. \\
&\quad + \sum_{j=1}^p \left[(a'_j - a''_j)^2 + \frac{1}{3}(a'_j - a''_j)[(\beta'_j + \beta''_j) - (\alpha'_j + \alpha''_j)] \right. \\
&\quad \left. \left. + \frac{1}{12}[(\alpha'_j + \alpha''_j)^2 + (\beta'_j + \beta''_j)^2] \right] \right) \\
&\quad + \sum_{i=1}^n (a_0 + \sum_{j=1}^p (a'_j - a''_j)x_{ij} - y_i)^2 + \frac{1}{3} \sum_{i=1}^n \left\{ (a_0 + \sum_{j=1}^p (a'_j - a''_j)x_{ij} - y_i) \right. \\
&\quad \times \left[(\beta_0 + \sum_{j=1}^p (a'_j\beta_{x_{ij}} + \beta'_jx_{ij} + a''_j\alpha_{x_{ij}} + \beta''_jx_{ij}) - \beta_{y_i}) \right. \\
&\quad \left. \left. - (\alpha_0 + \sum_{j=1}^p (a'_j\alpha_{x_{ij}} + \alpha'_jx_{ij} + a''_j\beta_{x_{ij}} + \alpha''_jx_{ij}) - \alpha_{y_i}) \right] \right\} \\
&\quad + \frac{1}{12} \sum_{i=1}^n \left[(\beta_0 + \sum_{j=1}^p (a'_j\beta_{x_{ij}} + \beta'_jx_{ij} + a''_j\alpha_{x_{ij}} + \beta''_jx_{ij}) - \beta_{y_i})^2 \right. \\
&\quad \left. + (\alpha_0 + \sum_{j=1}^p (a'_j\alpha_{x_{ij}} + \alpha'_jx_{ij} + a''_j\beta_{x_{ij}} + \alpha''_jx_{ij}) - \alpha_{y_i})^2 \right],
\end{aligned}$$

subject to

$$\begin{aligned}
a'_j &\geq \alpha'_j, a''_j \geq \beta''_j, \quad j = 1, \dots, p, \\
a_0 &\in \mathbb{R}, \alpha_0, \beta_0 \geq 0, \\
a'_j, a''_j, \alpha'_j, \alpha''_j, \beta'_j, \beta''_j &\geq 0, \quad j = 1, \dots, p.
\end{aligned}$$

To obtain the best choice of k and coefficients, we minimize the optimization problem (17) with Mathematica 11. Notice that to recognize the collinearity between fuzzy input variables, generalized variance inflation factor (VIF_G) introduced as follow

1.4. (5) For each fuzzy input variables

$$\text{VIF}_{G_i} = \frac{1}{1 - R_{G_j}^2}, \quad j = 1, 2, \dots, p, \tag{21}$$

where $R_{G_j}^2 = \frac{\sum_{i=1}^n d^2(\hat{X}_{ij}, \bar{X}_j)}{\sum_{i=1}^n d^2(X_{ij}, \bar{X}_j)}$ ($j = 1, 2, \dots, p$) is the coefficient of determination from the regression X_j on the other independent variables.

2.4. To evaluate the goodness of fit between the observed value $\tilde{y}_i$ and the estimated values $\tilde{Y}_i$, we define the mean square predict error (MPE) as follow

$$\text{MPE} = \frac{1}{n} \sum_{i=1}^n d^2(\tilde{y}_i, \tilde{Y}_i). \quad (22)$$

5 Performance Evaluation

5.1 Error of estimation

To evaluate the performance of fuzzy regression model, (10) introduced the following scaler as the difference between two membership functions $\tilde{Y}_i(x)$ and $\tilde{y}_i(x)$

$$\mathbb{E}_i = \int_{S_{\tilde{Y}_i} \cup S_{\tilde{y}_i}} |Y_i(x) - y_i(x)| dx, \quad i = 1, 2, \dots, n. \quad (23)$$

Where $S_{\tilde{Y}_i}$ and $S_{\tilde{y}_i}$ are the supports of TFNs $\tilde{Y}_i$ and $\tilde{y}_i$, respectively. Note that $\mathbb{E}_i$ is the non-overlapped area between the graphs of $\tilde{Y}_i(x)$ and $\tilde{y}_i(x)$, It is clear that the smaller value of $\mathbb{E}_i$ denotes the better estimation or better fitness of estimated response to the observed response. Let $\mathbb{E} = \sum_{i=1}^n \mathbb{E}_i$ denotes the total errors of estimation.

5.2 λ -Parameter

(6; 7) introduced a parameter named λ , which measures the total difference between the center of observed and estimated responses as follow

$$\lambda = \sum_{i=1}^n \lambda_i, \quad \text{where} \quad \lambda_i = \left| m_{y_i} - \sum_{j=0}^p m_{A_j} m_{X_{ij}} \right|. \quad (24)$$

5.3 Cross-validation

One of the other criteria for evaluating the performance of the model is cross-validation (cv). In order to measure the prediction of the model (see (16)), each time, the i th ($i = 1, 2, \dots, n$) observation is kept out for the validation set, then the $n - 1$ remaining observation are used to fit the model, and the fitted model is used to predict $\tilde{Y}_{(-i)}$ for the excluded observation. eventually, to obtain the distance between the observation $\tilde{y}_i$ and the predict $\tilde{Y}_{(-i)}$ (the predicted value with the $n - 1$ remaining observations), we calculate the weighted distance (11) between them as follow

$$\text{MPE}_{cv} = \frac{1}{n} \sum_{i=1}^n d^2(\tilde{y}_i, \tilde{Y}_{(-i)}). \quad (25)$$

6 Illustrative Example

We consider a dataset on Portland cement originally due to (15). These data come from an experimental investigation of the heat evolved during the setting and hardening of Portland cement of varied composition and the dependence of this heat on the percentages of four compounds in the clinkers from which the cement was produced. the four compounds are tricalcium aluminate ($3\text{CaO}.\text{Al}_2\text{O}_3$), tricalcium silicate ($3\text{CaO}.\text{SiO}_2$), tetracalcium aluminoferrite ($4\text{CaO}.\text{Al}_2\text{O}_3.\text{Fe}_2\text{O}_3$), and dicalcium silicate ($2\text{CaO}.\text{SiO}_2$), Which we will denote by X_1 , X_2 , X_3 , and X_4 , respectively. the heat evolved after 180 days of curing, which we will denote by Y , is measured in

Table 1: *Portland Cement Data*

i	$\tilde{X}_1$	$\tilde{X}_2$	$\tilde{X}_3$	$\tilde{X}_4$	$\tilde{Y}$
1	(7,0.77,0.91)	(26,2.86,3.38)	(6,0.66,0.78)	(60,6.6,7.8)	(78.5,8.635,10.20)
2	(1,0.12,0.2)	(29,3.48,5.8)	(15,1.8,3)	(52,6.24,10.4)	(74.3,8.91,14.86)
3	(11,2.09,1.21)	(56,10.64,6.16)	(8,1.52,0.88)	(20,3.8,2.2)	(104.3,19.81,11.47)
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
13	(10,1.7,1.1)	(68,11.56,7.48)	(8,1.36,0.88)	(12,2.04,1.32)	(109.4,18.59,12.03)

calories per gram of cement. We assemble the data in Table 1. values of the dependent variable, Y_i and independent variables X_{i1} , X_{i2} , X_{i3} , and X_{i4} were fuzzified(5) using an asymmetric triangular fuzzy number which its right and left spread proportionally to Y_i and X_{i1} , X_{i2} , X_{i3} , X_{i4} . Since the input variables are fuzzy, we use VIF_G to recognize multicollinearity between explanatory variables, we have

$$VIF_{G_1} = 3.705, \quad VIF_{G_2} = 41.4263, \quad VIF_{G_3} = 4.2225, \quad VIF_{G_4} = 15.847.$$

That shows collinearity between variables($\exists i = 1, 2, 3, 4, VIF_{G_i} \geq 5$), so we should use proposed model to calculate $\hat{\beta}_{Ridge}$. To obtain the best fit, we minimize the equation (17) under condition (18), with Mathematica 11. we have found that for $k = 1.3$, $MPE_{cv} = 7.2$ as for, $k = 0$ (ordinary fuzzy method) $MPE_{cv} = 7.71$, that is shown in presence of multicollinearity, the proposed model works very well. So we set $k = 1.3$ and calculate fuzzy ridge regression model as follow

$$\begin{aligned} \tilde{Y} &= (0.038, 0, 0.1) \oplus (2.167, 0.0003, 0) \otimes \tilde{X}_1 \oplus (1.16, 0.00003, 0) \otimes \tilde{X}_2 \\ &\oplus (0.729, 0.0004, 0) \otimes \tilde{X}_3 \oplus (0/49, 0.0002, 0) \otimes \tilde{X}_4, \end{aligned} \quad (26)$$

and also we calculate for $k = 1.3$ that, $\mathbb{E} = 38.29$ and $\lambda = 20.09$ to compare with $k = 0$ ($\mathbb{E} = 37.74$ and $\lambda = 19.82$), are almost the same.

7 Conclusion

In this paper, we proposed a new fuzzy ridge regression model with multiple fuzzy inputs and triangular fuzzy output. we realized that we can use this method successfully when there happens multicollinearity in the input variables or a linear model is inappropriate.

References

- [1] Balasundaram S. and Kapil (2011), Weighted fuzzy ridge regression analysis with crisp inputs and triangular fuzzy outputs, *Int. J. Advanced Intelligence Paradigms*, 3(1), 67-81.
- [2] Chang Y. H. O. (2001), Hybrid fuzzy least-squares regression analysis and its preliability measures, *Fuzzy Sets and Systems*, 119(2), 225-246.
- [3] Dubois D. and Prade H. (1980), *Fuzzy sets and systems: Theory and applications*, Academic Press, New York.
- [4] Draper N. R. and Smith H. (1980), *Applied Regression Analysis*, John Wiley, New York.
- [5] Farrokhi M. (2017), *Ridge Regression approach in Fuzzy Environment*, MSc Dissertation, Shahrood University of Technology, (in persian).

- [6] Hassanpour H. Maleki H. R. and Yaghoobi M. A. (2009b), A note on evaluation of fuzzy linear regression models by comparing membership functions, *Iranian Journal of Fuzzy Systems*, 6(2), 1-6.
- [7] Hassanpour H. Maleki H. R. and Yaghoobi M. A. (2010), Fuzzy linear regression model with crisp coefficients: a goal programming approach, *Iranian Journal of Fuzzy Systems*, 7(2), 19-39.
- [8] Hassanpour H. Maleki R. and Yaghoobi M. A. (2011), A goal programming approach to fuzzy linear regression with fuzzy input–output data, *Soft Computing* 15(8), 1569 -1580.
- [9] Hoerl A. E. and Kennard R. W. (1970), Ridge regression biased estimation for nonorthogonal problems, *Technometrics*, 12, 55-67.
- [10] Kim B. and Bishu R. R. (1998), Evaluation of fuzzy linear regression models by comparing membership functions, *Fuzzy Sets and Systems*, 100, 343-352.
- [11] Leinfellner W. and Eberlein G. (1992), *Fuzzy data analysis*, John Wiley, New York.
- [12] Nguyen H. T. and Walker E. A. (2005), *A first course in fuzzy logic*, Taylor & Francis, New York.
- [13] Sakawa M. (1993), *Fuzzy sets and interactive multiobjective optimization*, Plenum, New York.
- [14] Tanaka H. and Uejima S. and Asia K. (1982), Linear regression analysis with Fuzzy model, *IEEE Transaction Systems Man and Cybermatics*, 12(6), 903-907.
- [15] Woods H., Steinour H. H. and Starke H. R. (1932), Effect of composition of Portland cement on heat evolved during hardening, *Industrial and Engineering Chemistry*, 24(11), 1207-1214.
- [16] Wasserman L. (2006), *All of Nonparametric Statistics*, Springer Texts in Statistics. Springer , New York.
- [17] Xu R. and Li C. (2001), Multidimensional least squares fitting with a fuzzy model, *Fuzzy Sets and Systems*, 119, 215-223.
- [18] Zadeh L. A. (1978), Fuzzy sets as a basis for a theory of possibility, *Fuzzy Sets and Systems*, 1, 3-28.
- [19] Zimmermann H. J. (2001), *Fuzzy Set Theory and its Application*, Springer, New York.

A Fuzzy optimal planned replacement time based on availability and cost functions

Fatemeh Safaei¹ , Jafar Ahmadi, Bahram Sadeghpour Gildeh

Department of Statistics, Ferdowsi University of Mashhad

Abstract

Since the population or population's parameters of a system or subsystem are generally unknown, if we use the point estimates of parameters to estimate other characteristics of systems using a random sample from data in the past, these estimates of the system may fluctuate around the point estimates during a time interval. It follows that to use the point estimates is not suitable for the real cases. Therefore, it is more desirable to use the statistical confidence interval. Transferring the statistical confidence interval into the triangular fuzzy number gives a good sense to the values which belong to it. In this paper, we consider the parameters of lifetime distributions as triangular fuzzy numbers. Through these parameters, we calculate a fuzzy optimal planned replacement time based on cost and availability functions.

Keywords: Availability function, Cost function, Fuzzy number, Optimal planned replacement.

1 Introduction

The majority systems deteriorate with age and usage are influenced by stochastic failures during operation. Failures of systems incur high costs. Some of the failures are repairable and some of the failures need to be replaced. Finding an optimal replacement policy for a system becomes a major problem in reliability studies. Repair and replacement models have been extensively studied, and such models have been applied widely to nuclear power plant, electronics, aircraft industry, machinery and communications equipment, and industrial, environmental, military and medical systems. (4) presented the traditional age-replacement maintenance policy which a system is replaced at failure time or at age T , whichever occurs first. Several extensions of this policy have been investigated by researchers, e.g., (17) and (20). Furthermore, (6) considered the case of periodic replacement at times kT ($k = 1, 2, \dots$) and minimal

¹Speaker: safaei_fa123@yahoo.com

repair if the system fails otherwise. This model has been extended by (7) and (19). (21) introduced optimal preventive maintenance and repair policies for multi-state systems. (25) developed a maintenance policy with preventive repair and two types of failures. (14) proposed a bivariate replacement policy (n, T) for a cumulative shock damage process under the cumulative repair cost limit. Optimization problems of the cost and reliability-based maintenance strategies have been discussed in the literature, see for example, (11) and (13). However, only a few studies that have determined age replacement times based on the costs and the availability. (2) considered the case where the preventive maintenance epoch is randomized and obtained the optimal rate of preventive maintenance to maximize the system availability. (22) considered a case study on railway carriers and optimized inspection intervals reliability and availability. (16) provided some theoretical analyses of various models including classical replacement, preventive maintenance, and inspection policies. (1) proposed a case study maximized the availability and minimized the maintenance cost by using a multi-objective genetic algorithm. (23) considered a simple repairable system with delayed repair and studied the optimal replacement policy based on system age. (18) discussed availability and maintenance policy for systems subject to multiple failure modes. Their optimal maintenance policy maximizing steady-state availability and minimizing long-run average cost rate. (24) collected recent results and proposed some new models in age replacement policies.

In real-world situation, because of incomplete or non-obtainable information, data are often not so deterministic and the majority of these data can be assessed by human perception and human judgement. Therefore, they usually are fuzzy imprecise and so fuzzy theory can be applied in this problem to analyze qualitative verbal assessments. Fuzzy linguistic models permit the translation of verbal expression into numerical terms, thereby dealing quantitatively with the expression of the importance of various objectives.

(26) ever explained why fuzzy methodology can be used in operational research. He had the following three observations.

(1). Vague phenomenon: vague relations in the modeling of problems. That is, the problems themselves may be fuzzy in nature. For example, optimization problems may be attached with fuzzy objectives and/or fuzzy constraints.

(2). Informational vagueness: That is, the input data available to problems may be fuzzy. In particular, expert judgments are usually expressed in terms of linguistic variables and of fuzziness.

(3). Heuristic algorithms: Although the problems may be accurate, it can be too complex or expensive to get exact solution to them. In these circumstances heuristic or fuzzy algorithms can help.

The use of fuzzy methodology in system failure engineering can be traced back to (12). They then introduced the notion of component possibility as a reliability index to replace the notion of component probability. The main work of fuzzy methodology in system failure appeared in 1980s and after. Now fuzzy methodology has been widely applied in system failure engineering, e.g., in human reliability, hardware reliability, software reliability, structural reliability, maintenance and so on.

Fuzzy theory can be considered in the maintenance of systems. (5) proposed new approach for selecting optimum maintenance strategy using qualitative and quantitative data through interaction with the maintenance experts. This approach has been based on linear assignment method with some modifications to develop interactive fuzzy linear assignment method. (9) proposed a street-lighting lamps replacement problem and find a optimal fuzzy strategy for replacement the system.

In this paper, we assume that the system lifetime has a fuzzy parameters and we find fuzzy optimal planned replacement time based on cost and availability functions.

The rest of this paper is organized as follows. Section 2 contains the model assumptions and description. The optimization problem based on cost and availability functions are studied in Section 3. Some concluding remarks are presented in Section 4.

2 Model assumptions and optimization

Consider a system such that its lifetime is X with density and distribution functions $f_X(\cdot)$ and $F_X(\cdot)$, respectively; and suppose the system is maintained according to an age replacement maintenance policy. This age replacement policy in which a unit is replaced at constant time T after its installation or at failure, whichever occurs first. We call a specified time T the planned replacement time (or preventive replacement time) which ranges over $(0, \infty]$. The event $T = \infty$ represents that no replacement is made at all. Replacement at age X is non-planned and replacement at age T is planned replacement. It is assumed that failures are instantly detected and replaced with a new one, where its replacement time is negligible. We try to find time for preventive replacement according to the cost and availability functions and study the effect of the system parameters on the optimal planned replacement time.

Let X_T be the time to replacement of the system, then from the assumptions, we have $X_T = \min\{X, T\}$. The survival function of X_T is given by

$$\bar{F}_{X_T}(u) = P(X_T > u) = \bar{F}_X(u); \quad T \geq u. \quad (1)$$

As a result, from (1), the mean time to the replacement of the system can be expressed as

$$E(X_T) = \int_0^T \bar{F}_{X_T}(u) du. \quad (2)$$

To find optimal planned replacement time, we consider two common criteria namely, the costs and the availability of the system where the cost function should be minimized and the availability function should be maximized.

2.1 Cost function

Let c_F be the replacement cost at failure and c_T ($c_T < c_F$) be the replacement cost at time T . Then, the expected cost rate is

$$C(T) = \frac{c_F F_X(T) + c_T \bar{F}_X(T)}{\int_0^T \bar{F}_X(t) dt} = \frac{c_F - (c_T + c_F) \bar{F}_X(T)}{\int_0^T \bar{F}_X(t) dt}. \quad (3)$$

due to $\bar{F}_X(T) = 1 - \int_0^T \lambda(u) \bar{F}_X(u) du$, we have

$$C(T) = \frac{c_T - (c_T - c_F) \int_0^T \lambda(u) \bar{F}_X(u) du}{\int_0^T \bar{F}_X(t) dt}, \quad (4)$$

where $\lambda(t) = \frac{f(t)}{\bar{F}_X(t)}$.

We want to find T_C so that $C(T)$ needs to be minimized at $T = T_C$. (4) and (24) presented If $\lambda(t)$ increases strictly with t and optimal T_C ($0 < T_C < \infty$) to minimize $C(T)$ satisfies

$$\lambda(T) \int_0^T \bar{F}_X(u) du - F(T) = \frac{c_T}{c_F - c_T}. \quad (5)$$

Let

$$M(T) = \int_0^T (\lambda(T) - \lambda(u)) \bar{F}_{X_T}(u) du + \frac{c_T}{c_T - c_F}. \quad (6)$$

It is deduced that $\frac{d}{dT} C(T)$ equals 0, is greater than 0, or is less than 0, if $M(T)$ equals 0, is greater than 0, or is less than 0, respectively.

As a result, if an optimal T exist, then it satisfies $M(T_C) = 0$.

Suppose X has a Weibull distribution for every parameters $\tilde{\lambda}$ and $\tilde{k}$ with density function $f_X(x) = \frac{\tilde{k}}{\tilde{\lambda}} \left(\frac{x}{\tilde{\lambda}}\right)^{\tilde{k}-1} e^{-\left(\frac{x}{\tilde{\lambda}}\right)^{\tilde{k}}}$ and distribution function $F_X(x) = 1 - e^{-\left(\frac{x}{\tilde{\lambda}}\right)^{\tilde{k}}}$ where $\tilde{\lambda}$ and $\tilde{k}$ are fuzzy number. Therefore we have

$$C(\tilde{T}) = \frac{c_F + (c_T - c_F) e^{-\left(\frac{\tilde{T}}{\tilde{\lambda}}\right)^{\tilde{k}}}}{\int_0^{\tilde{T}} e^{-\left(\frac{u}{\tilde{\lambda}}\right)^{\tilde{k}}} du}, \quad (7)$$

and

$$M(\tilde{T}) = \frac{\tilde{k}}{\tilde{\lambda}^{\tilde{k}}} \int_0^{\tilde{T}} (\tilde{T}^{\tilde{k}-1} - u^{\tilde{k}-1}) e^{-\left(\frac{u}{\tilde{\lambda}}\right)^{\tilde{k}}} du - \frac{c_T}{c_T - c_F}. \quad (8)$$

For $0 \leq \alpha \leq 1$, considering Buckley approach ((8)), we have

$$\tilde{T}_C[\alpha] = [\tilde{T}_C^-, \tilde{T}_C^+], \quad (9)$$

where

$$\tilde{T}_C^-[\alpha] = \min\{T_C | T_C = \arg \min C(T); \lambda \in \tilde{\lambda}[\alpha], k \in \tilde{k}[\alpha]\}, \quad (10)$$

and

$$\tilde{T}_C^+[\alpha] = \max\{T_C | T_C = \arg \min C(T); \lambda \in \tilde{\lambda}[\alpha], k \in \tilde{k}[\alpha]\}. \quad (11)$$

Suppose $\tilde{\lambda} = (3, 0.1, 0.1)_T$ and $\tilde{k} = (5, 0.1, 0.1)_T$ are two triangular fuzzy number and $c_F = 20$ and $c_T = 2$, then the fuzzy optimal planned replacement time T , is plotted in Figure 1.

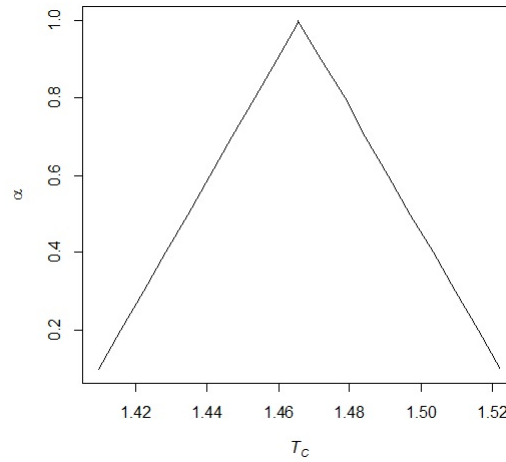


Figure 1: Fuzzy optimal planned replacement time.

2.2 Availability function

In the maintenance and reliability theory, the availability of a system is defined in many ways; for example: pointwise availability, interval availability, limiting interval availability, multiple cycle availability, see for example, (15) (pages 9-10). In this paper, we consider the limiting interval availability which is given by

$$AV = \frac{MTTF}{MTTF + MTTR}, \quad (12)$$

where MTTF (mean time to failure) is the length of time that the system is expected to last in operation (mean time which the system is up) and MTTR represents the expected time required to replace the system (mean time which the system is down). Both of them are nonnegative, then $0 < AV < 1$. Notice that AV sometimes called simply availability.

Consider the unit undergoes repair after failure according to a general distribution $G_F(t)$ with finite mean μ_F , and the unit can be quickly resumed to normal operation after repairs. In addition, the unit is replaced at a planned time T according to a general distribution $G_T(t)$ with finite mean μ_T ($\mu_T < \mu_F$). Then, the availability is

$$\begin{aligned} AV(T) &= \frac{\int_0^T \bar{F}_X(u) du}{\int_0^T \bar{F}_X(u) du + \mu_F(1 - \bar{F}_X(T)) + \mu_T \bar{F}_X(T)} \\ &= \frac{\int_0^T \bar{F}_X(u) du}{\int_0^T \bar{F}_X(u) du + (\mu_T - \mu_F) \bar{F}_X(T) + \mu_F}, \end{aligned} \quad (13)$$

due to $\bar{F}_X(T) = 1 - \int_0^T \lambda(u) \bar{F}_X(u) du$, we have

$$AV(T) = \frac{\int_0^T \bar{F}_X(u) du}{\int_0^T (1 - (\mu_F - \mu_T) \lambda(u)) \bar{F}_X(u) du + \mu_T}. \quad (14)$$

In order to find the optimal T , say T_A^* , so that $AV(T)$ needs to be maximized at $T = T_A^*$. From (13) we derive

$$\frac{d}{dT} AV(T) = \frac{\bar{F}_{X_T}(T) \left\{ -(\mu_F - \mu_T) \int_0^T (\lambda(T) - \lambda(u)) \bar{F}_{X_T}(u) du + \mu_T \right\}}{\left(\int_0^T (1 + (\mu_F - \mu_T) \lambda(u)) \bar{F}_{X_T}(u) du + \mu_T \right)^2},$$

since, by the assumptions, $\bar{F}_{X_T}(u)$ and $\lambda(u)$ are continuous functions. Upon taking

$$Q(T) = \int_0^T (\lambda(T) - \lambda(u)) \bar{F}_{X_T}(u) du + \frac{\mu_T}{\mu_T - \mu_F}, \quad (15)$$

it is deduced that $\frac{d}{dT} AV(T)$ equals 0, is greater than 0, or is less than 0, if $Q(T)$ equals 0, is less than 0, or is greater than 0, respectively. As a result, if an optimal T_A^* exist, then it satisfies $Q(T_A^*) = 0$. On the other hand

$$\lim_{T \rightarrow 0} Q(T) = \frac{\mu_T}{\mu_T - \mu_F} < 0, \quad (16)$$

and

$$\frac{d}{dT} Q(T) = \int_0^T \bar{F}_{X_T}(u) du \frac{d}{dT} \lambda(T). \quad (17)$$

If $\lambda(T)$ is an increasing function of T (IFR), then from (17), $Q(T)$ is increasing in T . By (15), $Q(T)$ can be reexpressed as

$$Q(T) = \left(1 + (\mu_F - \mu_T) \lambda(T) - \frac{1}{AV(T)} \right) \frac{\int_0^T \bar{F}_{X_T}(u) du}{\mu_F - \mu_T}. \quad (18)$$

Summing up, by using (18), we can verify the existence and uniqueness of T_A^* which is presented in the next proposition. Let $AV(T)$ and $Q(T)$ be as given in (13) and (18), respectively. Then:

If $AV(\infty) > \frac{1}{1 + (\mu_F - \mu_T) \lambda(\infty)}$, then $Q(\infty) > 0$ and there exist finite T_A^* that $\frac{d}{dT} AV(T)|_{T=T_A^*} = 0$ and T_A^* maximizing $AV(T)$.

If $L(u)$ is strictly increasing in u , then $Q(T)$ is strictly increasing in T , hence there exists a unique and finite maximum T_A^* provided that $\frac{1}{AV(\infty)} < 1 + (\mu_F - \mu_T) \lambda(\infty)$.

If $L(u)$ is non-decreasing in u , then $Q(T)$ is non-decreasing in T , and $T = \infty$ is optimal provided that $\frac{1}{AV(\infty)} > 1 + (\mu_F - \mu_T)\lambda(\infty)$.

It may be noted that by equation (15), $Q(T)$ is an increasing function of μ_F and a decreasing function of μ_T . Consequently, with increasing μ_F , the planned replacement needs to be done at an earlier time and with increasing μ_T , the planned replacement needs to be done at a later time.

Suppose X has a Weibull distribution for every parameters $\tilde{\lambda}$ and $\tilde{k}$ with density function $f_X(x) = \frac{\tilde{k}}{\tilde{\lambda}} \left(\frac{x}{\tilde{\lambda}}\right)^{\tilde{k}-1} e^{-\left(\frac{x}{\tilde{\lambda}}\right)^{\tilde{k}}}$ and distribution function $F_X(x) = 1 - e^{-\left(\frac{x}{\tilde{\lambda}}\right)^{\tilde{k}}}$. Therefore we have

$$AV(\tilde{T}) = \frac{\int_0^{\tilde{T}} e^{-\left(\frac{u}{\tilde{\lambda}}\right)^{\tilde{k}}} du}{\mu_F + (\mu_T - \mu_F)e^{-\left(\frac{\tilde{T}}{\tilde{\lambda}}\right)^{\tilde{k}}} + \int_0^{\tilde{T}} e^{-\left(\frac{u}{\tilde{\lambda}}\right)^{\tilde{k}}} du}, \quad (19)$$

and

$$Q(\tilde{T}) = \frac{\tilde{k}}{\tilde{\lambda}^{\tilde{k}}} \int_0^{\tilde{T}} (\tilde{T}^{\tilde{k}-1} - u^{\tilde{k}-1}) e^{-\left(\frac{u}{\tilde{\lambda}}\right)^{\tilde{k}}} du - \frac{\mu_T}{\mu_T - \mu_F}. \quad (20)$$

For $0 \leq \alpha \leq 1$, considering Buckley approach ((8)), we have

$$\tilde{T}_A[\alpha] = [\tilde{T}_A^-, \tilde{T}_A^+], \quad (21)$$

where

$$\tilde{T}_A^-[\alpha] = \min\{T_A | T_A = \arg \max AV(T); \lambda \in \tilde{\lambda}[\alpha], k \in \tilde{k}[\alpha]\}, \quad (22)$$

and

$$\tilde{T}_A^+[\alpha] = \max\{T_A | T_A = \arg \max AV(T); \lambda \in \tilde{\lambda}[\alpha], k \in \tilde{k}[\alpha]\}. \quad (23)$$

Suppose $\tilde{\lambda} = (3, 0.2, 0.2)_T$ and $\tilde{k} = (3, 0.2, 0.2)_T$ are two triangular fuzzy number and $\mu_F = 2$ and $\mu_T = 1$, then the fuzzy optimal planned replacement time T , is plotted in Figure 2.

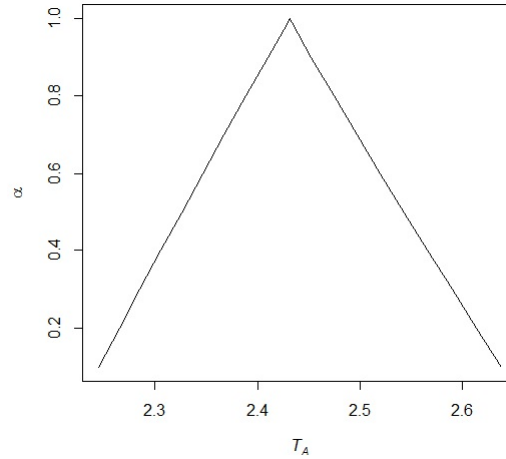


Figure 2: Fuzzy optimal planned replacement time.

References

- [1] Adhikary, D. D., Bose, G. K., Jana, D. K., Bose, D. and Mitra, S. (2016). Availability and cost-centered preventive maintenance scheduling of continuous operating series systems using multi-objective genetic algorithm: A case study. *Quality Engineering*, 28, 352-357.

- [2] Angus, J. E., Yin, M. L. and Trivedi, K. (2012). Optimal random age replacement for availability. *International Journal of Reliability, Quality and Safety Engineering*, 19, doi.org/10.1142/S0218539312500210.
- [3] Aven, T. (1992). *Reliability and Risk Analysis*, Elsevier Applied Science, London.
- [4] Barlow, R. E. and Proschan, F. (1965). *Mathematical Theory of Reliability*, John Wiley and Sons, New York.
- [5] Bashiri, M., Badri, H. and Hejazi, T. H. (2011). Selecting optimum maintenance strategy by fuzzy interactive linear assignment method. *Applied Mathematical Modelling*, 35, 152-164.
- [6] Boland, P. J. and Proschan, F. (1982). Periodic replacement with increasing minimal repair costs at failure. *Operations Research*, 30, 1183-1189.
- [7] Boland, P.J. and Proschan, F. (1983). The reliability of k out of n systems. *Annals of Probability*, 11, 760-764.
- [8] Buckley, J.J. (2006). *Fuzzy Probability and Statistics*. Springer, Heidelberg.
- [9] Cai, K. Y. (1996). *Introduction to Fuzzy Reliability*, Kluwer Academic Publishers, Dordrecht.
- [10] Chen, Y.L., (2012), A bivariate optimal imperfect preventive maintenance policy for a used system with two-type shocks, *Computers and Industrial Engineering*, 63, 1227-1234.
- [11] Garbatov, Y. and Guedes S. C. (2001). Cost and reliability based strategies for fatigue maintenance planning of floating structure. *Reliability Engineering and System Safety*, 73, 293-301.
- [12] Kaufmann, A. and Bonaert, A. P. (1977). Introduction to the Theory of Fuzzy Subsets-vol. 1: Fundamental Theoretical Elements, *IEEE Transactions on Systems*, 7, 495-496.
- [13] Lapa C. M. F., Pereira C. M. N. A. and de Barros M. P. (2006). A model for preventive maintenance planning by genetic algorithms based in cost and reliability, *Reliability Engineering and System Safety*, 91, 233-240.
- [14] Lai M., Chen C. and Hariguna T. (2017). A bivariate optimal replacement policy with cumulative repair cost limit for a two-unit system under shock damage interaction, *Brazilian Journal of Probability and Statistics*, 2, 353-372.
- [15] Nakagawa, T. (2005). *Maintenance Theory of Reliability*, London, Springer.
- [16] Nakagawa, T. (2014). *Random Maintenance Policies*, London, Springer.
- [17] Nakagawa, T. and Kowada, M. (1983). Analysis of a system with minimal repair and its application to replacement policy. *European Journal of Operation Research*, 12, 176-182.
- [18] Qiu, Q., Cui, L. and Gao, H. (2017). Availability and maintenance modelling for systems subject to multiple failure modes. *Computers and Industrial Engineering*, 108, 192-198.
- [19] Sheu, S. H. (1996). A modified block replacement policy with two variables general and random minimal repair cost. *Journal of Applied Probability*, 33, 557-572.
- [20] Sheu, S., Chang, C. C. (2009). An extended periodic preventive maintenance model with age-dependent failure type. *IEEE Transformation on Reliability*, 85, 397-405.
- [21] Sheu, S., Chang, C. C., Chen, Y. and Zhang, Z. (2015). Optimal preventive maintenance and repair policies for multi-state systems. *Reliability Engineering and System Safety*, 40, 78-87.
- [22] Wolde, M. T. and Ghobbar, A. A. (2013). Optimizing inspection intervals-reliability and availability in terms of a cost model:A case study on railway carriers. *Reliability Engineering and System Safety*, 114, 137-147.
- [23] Zhang, Y. L. and Wang, G. J. (2017). An optimal age-replacement policy for a simple repairable system with delayed repair. *Communications in Statistics-Theory and Methods*, 46, 2837-2850.
- [24] Zhao, X. Al-Khalifa, K. N., Hamouda, A. M. and Nakagawa, T. (2017) Age replacement models: A summary with new perspectives and methods. *Reliability Engineering and System Safety*, 161, 95-105.
- [25] Zheng, Z., Zhou, W., Zheng, Y. and Wu, Y. (2016). Optimal maintenance policy for a system with preventive repair and two types of failures. *Computers and Industrial Engineering*, 98, 102-112.
- [26] Zimmermann, H. J. (1983). Using Fuzzy Sets in Operational Research, *European Journal of Operational Research*, 3, 201-216.

Proceeding of

The 8th Seminar on

Fuzzy Statistics and Probability

Department of statistics,
Ferdowsi University of Mashhad,
Mashhad, Iran.

9-10 May, 2018



Ferdowsi University
of Mashhad



Proceeding of 8th Seminar on Fuzzy Statistics and Probability

May 9 -10, 2018

In memory of
Prof. N. Arghami

