



مجموعه مقالات

هفتمین سمینار آمار و احتمال فازی

دانشگاه بیرجند

۱۳ و ۱۴ اردیبهشت ۱۳۹۶



اعضای کمیته علمی

دکتر سید محمود طاهری	دانشگاه تهران
دکتر بهرام صادق پور گیلده	دانشگاه فردوسی مشهد
دکتر غلامرضا محتشمی برزادران	دانشگاه فردوسی مشهد
دکتر عباس پرچمی	دانشگاه شهید باهنر کرمان
دکتر محمدرضا ربیعی	دانشگاه صنعتی شاهرود
دکتر جلال چاچی	دانشگاه سمنان
دکتر رضا زارعی	دانشگاه گیلان
دکتر حسن حسن پور	دانشگاه بیرجند
دکتر محمدقاسم اکبری	دانشگاه بیرجند

کمیته برگزار کننده

ریاست دانشگاه بیرجند
معاونت پژوهشی و فناوری دانشگاه بیرجند
معاونت اداری و مالی دانشگاه بیرجند
معاونت فرهنگی و اجتماعی دانشگاه بیرجند
معاونت دانشجویی دانشگاه بیرجند
معاونت آموزشی و تحصیلات تکمیلی دانشگاه بیرجند
روابط عمومی دانشگاه بیرجند
نقلیه دانشگاه بیرجند
دانشکده علوم ریاضی و آمار دانشگاه بیرجند
گروه آمار دانشگاه بیرجند
کمیته علمی سمینار
کمیته اجرایی سمینار
دبیرخانه سمینار

اعضای کمیته اجرایی

رئیس دانشکده علوم ریاضی و آمار	دکتر محمد مهدی نصرآبادی
معاون آموزشی دانشکده علوم ریاضی و آمار	دکتر حاجی محمد محمدی نژاد
مسئول پژوهشی دانشکده علوم ریاضی و آمار	دکتر مجید رضایی
مدیر گروه آمار	دکتر جواد اطمینان
دبیر کمیته اجرایی	دکتر محسن عارفی
عضو کمیته اجرایی	دکتر محمد خراشادی زاده
عضو کمیته اجرایی	دکتر هادی علیزاده نوقابی
عضو کمیته اجرایی	دکتر مجید چهکندی
مسئول طراحی	حمید سیف الهی مقدم
همکار دبیرخانه	الهه پامرغی
همکار دبیرخانه	رضا عابدی پور
همکار دبیرخانه	حمیدرضا ظهور سلیمانی
همکار دبیرخانه	عاطفه احمدی
همکار دبیرخانه	هدیه افتخاری
همکار دبیرخانه	مریم مالکی

مدعوین

۱. دکتر سید محمود طاهری، عضو هیات علمی دانشگاه تهران
۲. دکتر بهرام صادق‌پور گیلده، عضو هیات علمی دانشگاه فردوسی مشهد
(نماینده‌ی انجمن سیستم‌های فازی ایران)
۳. دکتر غلامرضا محتشمی برزادران، عضو هیات علمی دانشگاه فردوسی مشهد
(نماینده‌ی انجمن آمار ایران)

سخنرانان مهمان

۱. دکتر عباس پرچمی، عضو هیات علمی دانشگاه شهید باهنر کرمان
۲. دکتر محمدرضا ربیعی، عضو هیات علمی دانشگاه صنعتی شاهرود
۳. دکتر جلال چاچی، عضو هیات علمی دانشگاه سمنان
۴. دکتر رضا زارعی، عضو هیات علمی دانشگاه گیلان

فهرست

۸	تأثیر خطاهای بازرسی بر منحنی مشخصه عملکرد ... رباب افشاری، بهرام صادقیورگیلده
۱۷	کیفیت فازی: نمودارهای کنترل شوهارت عباس پرچمی
۲۵	رگرسیون وزنی فازی کمترین قدرمطلق خطا جلال چاچی، فائزه ترکیان
۳۲	کاربرد رگرسیون فازی در مدل‌بندی شاخص‌های کیفی آب جلال چاچی، مریم میرزایی، بهنام عرب حسن آبادی
۴۰	منطق فازی، احتمال، استنتاج منطقی و احتمالاتی سید محمد امین خاتمی
۴۸	نگاهی دیگر به معنی‌شناسی در منطق فازی سید محمد امین خاتمی، قاسم خاکشور
۵۶	معرفی شاخص‌های کارایی فازی چندمتغیره بنیتا دولت‌زاده، مصطفی گریوانی، ماشالله ماشین‌چی
۶۴	رگرسیون ریج فازی در حضور داده‌های پرت برای ضرایب و خروجی فازی محمد رضا ربیعی، معصومه فرخی، محمد آرشی
۷۳	تحلیل رگرسیون استوار فازی براساس کمترین قدرمطلق خطاها ... محمد رضا ربیعی، الهام رحیمیان، داود شاهسونی
۸۲	طراحی یک سیستم فازی-عصبی خود سازمانده با قابلیت یادگیری برخظ ... شیرین ریخته گرمشهد، حمید طباطبائی
۹۷	روش‌هایی در تحلیل قابلیت اطمینان سیستم‌های تعمیرپذیر در شرایط نادقیق رضا زارعی

- ۱۱۶ برآورد ریبج در مدل رگرسیون فازی با متغیرهای مستقل و پاسخ فازی
سیده راضیه سجادی، محسن عارفی، محمد خراشادی زاده
- ۱۲۳ انتخاب مدل مناسب در رگرسیون لوجستیک فازی
فاطمه سلمانی، سید محمود طاهری، علیرضا ابدی، حمید علوی مجد
- ۱۳۴ شاخص‌های کارایی فرایند تعمیم یافته با اعداد فازی و کیفیت فازی
نگار شیخی، محمدرضا ربیعی، عباس پرچی
- ۱۴۳ نمودار کنترل کیفیت فازی- آماری فرایند برای جمعیت چوله
مهسا صعصعی، ماشالله ماشین‌چی، رضا پور موسی، فرزانه هاشمی
- ۱۵۲ نقش خودتنظیمی یادگیری در عملکرد تحصیلی فراگیرندگان: فرا رگرسیون فازی
سید محمود طاهری، اصغر شیرعلی‌پور، مسعود اسدی
- ۱۶۷ آزمون فرضیه در محیط نادقیق با استفاده از روش بیز امکانی
ملیحه عابدینی
- ۱۷۶ استوارسازی تابع عضویت فازی فاصله مبنا برای ماشین بردار پشتیبان
ماندانا محمدی، مجید سرمد
- ۱۸۵ آزمون فرضیه‌های دقیق بر اساس داده‌های فازی نرمال با رویکرد p -مقدار
محبوبه سادات مدنی، محمدرضا ربیعی
- ۱۹۲ نمودار تعداد اقلام معیوب با استفاده از درجه عدم انطباق
میترا مهدی‌پور، محمدرضا ربیعی
- ۲۰۰ نمودارهای کنترل کیفیت فازی برای توزیع چوله نرمال
فرزانه هاشمی، ماشالله ماشین‌چی، احد جمالیزاده، مهسا صعصعی
- ۲۰۹ رگرسیون نیمه پارامتری فازی بر اساس خوشه‌بندی فازی
معصومه اسداللهی، محمدقاسم اکبری



تأثیر خطاهای بازرسی بر منحنی مشخصه عملکرد طرح نمونه‌گیری تعویقی چندگانه در محیط فازی

رباب افشاری^۱ * و بهرام صادقپور گیلده^۲

^۱robab.afshari@mail.um.ac.ir

^۲sadeghpour@um.ac.ir

چکیده. دقیق بودن نسبت اقلام معیوب انباشته و ایده‌آل بودن شرایط بازرسی، دو فرض معمول در طراحی هر طرح نمونه‌گیری برای پذیرش است. در مطالعه حاضر، نویسندگان ضمن پیشنهاد طرح نمونه‌گیری تعویقی چندگانه در محیط فازی، زمانیکه نسبت اقلام معیوب انباشته یک کمیت نادقیق باشد، به توسیع طرح پیشنهادی در حضور خطاهای بازرسی، می‌پردازند. نتایج به دست آمده نشان داد که قدرت عملکرد طرح پیشنهادی در تشخیص کیفیت خوب و بد انباشته‌ای از تولیدات، متأثر از خطاهای بازرسی بوده و در نتیجه منجر به تصمیم‌گیری غیرواقعی می‌شود.

۱. مقدمه

به منظور تصمیم‌گیری جهت پذیرش یا رد انباشته‌ای از تولیدات، یکی از ابزارهای موجود در کنترل کیفیت آماری طرح نمونه‌گیری برای پذیرش است. با نظر به اینکه مشخصه کیفی مورد نظر را بتوان به صورت یک کمیت عددی بیان کرد یا نه، طرح‌های نمونه‌گیری به ترتیب به دو دسته متغیر و وصفی تقسیم‌بندی می‌شوند.

طرح نمونه‌گیری تعویقی (وابسته به گذشته) چندگانه^۱ (MDS) که عضوی از خانواده طرح‌های نمونه‌گیری شرطی محسوب می‌شود اولین بار در مرجع [۲۳] مطرح شد. در خصوص طرح مذکور، که از اطلاعات موجود در انباشته‌های آتی (گذشته) در ساخت معیار تصمیم‌گیری استفاده می‌کند، مطالعات متعددی صورت گرفته است. در مرجع [۲۴] طرح MDS برای صفات اندازه‌پذیر، بر پایه توزیع نرمال زمانی که حدود مشخصات فنی دو طرفه باشد معرفی شد. طرح MDS به روش بیزی با توزیع پیشین گاما در [۱۴] بررسی شده است. جهت جزئیات بیشتر در خصوص طرح مذکور می‌توان به [۱۷]، [۱۸] و [۲۵] مراجعه کرد.

دو فرض معمول در طراحی هر نوع طرح نمونه‌گیری، دقیق بودن مقدار نسبت اقلام معیوب انباشته (p) و ایده‌آل بودن شرایط بازرسی است. اما در برخی موارد، امکان برآورد دقیقی از

واژگان کلیدی. کنترل کیفیت، طرح نمونه‌گیری تعویقی چندگانه، خطاهای بازرسی، حساب اعداد فازی.

* سخنران

^۱Multiple Deferred (Dependent) State Sampling Plan

نسبت اقلام معیوب انباشته وجود ندارد. در چنین مواقعی که نسبت مذکور به صورت یک کمیت زبانی عنوان می‌شود، روش کلاسیک کارایی نداشته و باید از نظریه مجموعه‌های فازی جهت طراحی طرح استفاده شود. در [۷] کلاسی از طرح‌های نمونه‌گیری ساده مبتنی بر بهینه‌سازی فازی مطرح شده است. به منظور مطالعات بیشتر در این خصوص می‌توان به [۱]، [۳]، [۴]، [۱۲] و [۱۹] - [۲۲] مراجعه کرد.

از سوی دیگر، در واقعیت امکان اینکه بازرسان در طول اجرای بازرسی، مرتکب خطا شوند وجود دارد. در حالت کلی، دو نوع خطای بازرسی وجود دارد: طبقه‌بندی کالای سالم به عنوان ناسالم و برعکس دسته‌بندی کالای معیوب به عنوان کالای سالم، که به ترتیب خطای نوع اول و دوم نامیده می‌شوند. در مرجع [۶] تاثیر خطاهای بازرسی بر متوسط کیفیت خروجی بررسی شده است. تاثیر خطاهای بازرسی بر روی طرح‌های نمونه‌گیری برای پذیرش براساس اطلاعات انباشته‌های پیرامون توسط [۱۶] مطالعه شد. در مرجع [۱۱] نویسندگان، طرح نمونه‌گیری برای پذیرش بیزی در حضور خطاهای بازرسی را بررسی کردند. جهت جزییات بیشتر می‌توان به [۲] و [۸] - [۱۰]، مراجعه کرد. در بخش بعد طرح نمونه‌گیری تعویقی چندگانه فازی با مشخصه کیفی وصفی و مقدار p نادقیق که به اختصار با نماد FMDS^۲ نشان داده می‌شود معرفی و طرح پیشنهادی در حضور خطای بازرسی، در بخش سوم مورد مطالعه قرار می‌گیرد. در بخش چهارم، طرح پیشنهادی در حالت خاص مورد بررسی قرار گرفته و نتایج اخذ شده در بخش پایانی آورده شده است.

۲. طرح نمونه‌گیری تعویقی چندگانه

در این بخش ابتدا به یادآوری طرح MDS در حالت کلاسیک پرداخته و در ادامه الگوریتم اجرای طرح مذکور در محیط فازی معرفی می‌شود.

۱.۲. الگوریتم اجرای طرح MDS در حالت کلاسیک.

گام اول: یک نمونه به حجم n از انباشته k ام برداشته و تعداد اقلام معیوب (d) در آن شمارش می‌شود.

گام دوم: اگر $d \leq c_1$ انباشته k ام را پذیرفته، اگر $d > c_2$ انباشته k ام رد می‌شود و در صورتی که $c_1 < d \leq c_2$ انباشته k ام در صورتی پذیرفته می‌شود که همه m انباشته متوالی بعدی پذیرفته شوند.

بنابراین طرح MDS توسط چهار پارامتر c_1 ، c_2 ، m و n مشخص شده و با نماد $MDS(c_1, c_2)$ نشان داده می‌شود. اگر متغیر تصادفی D نمایانگر تعداد اقلام معیوب در نمونه‌ای به حجم n و p نسبت اقلام معیوب در انباشته باشد در آن صورت با فرض بزرگ بودن حجم انباشته، متغیر تصادفی D دارای توزیع دوجمله‌ای با پارامترهای n و p خواهد بود. مطابق با [۲۳]، احتمال پذیرش انباشته k ام $(P_{a:k}(p))$ برابر است با:

$$P_{a:k}(p) = P(D \leq c_1) + P(c_1 < D \leq c_2)P_{a:k}^m(p),$$

$$= \sum_{d=0}^{c_1} P_d(p) + \left(\sum_{d=c_1+1}^{c_2} P_d(p) \right) P_{a:k}^m(p), \quad (۱.۲)$$

که در آن $P_d(p)$ برابر احتمال وجود دقیقاً d قلم کالای معیوب در نمونه‌ای به حجم n ، زمانی که نسبت اقلام معیوب انباشته p باشد، است.

۲.۲. الگوریتم اجرای طرح FMDS. فرض کنید مقدار نسبت اقلام معیوب انباشته به صورت عدد فازی مثلی $\tilde{p} = (a_1, a_2, a_3)$ باشد، آنگاه $-\lambda$ برش $\tilde{p}$ ، به ازای هر $\lambda \in [0, 1]$ برابر است با [۱]:

$$\tilde{p}[\lambda] = [a_1 + (a_2 - a_1)\lambda, a_3 - (a_3 - a_2)\lambda], \quad (۲.۲)$$

در آن صورت الگوریتم اجرای طرح MDS با پارامتر فازی $\tilde{p}$ یعنی $\text{FMDS}(c_1, c_2)$ ، دقیقاً مشابه طرح $\text{MDS}(c_1, c_2)$ بوده به جز اینکه به دلیل ابهام در مقدار نسبت اقلام معیوب، با فرض بزرگ بودن حجم انباشته، متغیر تصادفی D ، دارای توزیع دوجمله‌ای فازی با پارامترهای n و $\tilde{p}$ است [۵]. در این صورت $-\lambda$ برش احتمال پذیرش فازی انباشته k ام $(\tilde{P}_{a:k}(\tilde{p})[\lambda])$ برابر است با:

$$\begin{aligned} \tilde{P}_{a:k}(\tilde{p})[\lambda] &= \tilde{P}(D \leq c_1)[\lambda] + \tilde{P}(c_1 < D \leq c_2)[\lambda] \times \tilde{P}_{a:k}^m(\tilde{p})[\lambda] \\ &= \left(\sum_{d=0}^{c_1} \tilde{P}_d(\tilde{p}) \right) [\lambda] + \left(\sum_{d=c_1+1}^{c_2} \tilde{P}_d(\tilde{p}) \right) [\lambda] \times \tilde{P}_{a:k}^m(\tilde{p})[\lambda], \quad (۳.۲) \end{aligned}$$

که در آن $\tilde{P}_d(\tilde{p})$ برابر احتمال فازی وجود دقیقاً d کالای معیوب در نمونه‌ای به حجم n ، زمانی که نسبت اقلام معیوب انباشته برابر با $\tilde{p}$ باشد، است. همچنین مطابق روش باکلی داریم:

$$\begin{aligned} \tilde{P}(D \leq c_1)[\lambda] &= \left\{ \sum_{d=0}^{c_1} c_n^d p^d (1-p)^{n-d} : s \right\} \\ &= [P_I^L[\lambda], P_I^U[\lambda]], \\ \tilde{P}(c_1 < D \leq c_2)[\lambda] &= \left\{ \sum_{d=c_1+1}^{c_2} c_n^d p^d (1-p)^{n-d} : s \right\} \\ &= [P_{II}^L[\lambda], P_{II}^U[\lambda]], \end{aligned}$$

که در آن $P_I^L[\lambda]$ و $P_I^U[\lambda]$ به ترتیب می‌نیم و ماکزیم $\tilde{P}(D \leq c_1)[\lambda]$ هستند. می‌نیم و ماکزیم عبارت $\tilde{P}(c_1 < D \leq c_2)[\lambda]$ نیز به ترتیب برابر $P_{II}^L[\lambda]$ و $P_{II}^U[\lambda]$ می‌باشد. در تمامی عبارات فوق، s معادل عبارت $p \in \tilde{p}[\lambda]$ ، $q \in \tilde{q}[\lambda]$ و $p + q = 1$ است به طوریکه $\tilde{q}[\lambda] = \{1 - p ; p \in \tilde{p}[\lambda]\}$.

۳. طرح FMDS در حضور خطای بازرسی

در این بخش طرح FMDS را زمانیکه بازرسی در شرایط ایده‌آل انجام نشود، بررسی می‌کنیم. چنانچه اشاره شد در طول عمل بازرسی، بازرسان ممکن است مرتکب دو نوع خطای بازرسی، خطای نوع اول بازرسی (e_1) و نوع دوم بازرسی (e_2) شوند. برای این منظور پیشامدهای زیر را در نظر بگیرید:

E_1 : پیشامد آنکه یک محصول، واقعا معیوب باشد.

E_2 : پیشامد آنکه یک محصول، واقعا سالم باشد.

A : پیشامد آنکه یک کالا بعد از بازرسی به عنوان یک محصول معیوب در نظر گرفته شود.

$A|E_2$: پیشامد آنکه یک کالای سالم به عنوان یک کالای معیوب در نظر گرفته شود.

$A^c|E_1$: پیشامد آنکه یک کالای معیوب به عنوان یک کالای سالم، در نظر گرفته شود. (A^c متمم پیشامد A است)

فرض کنید $\tilde{p}$ و $\tilde{p}_e$ به ترتیب نسبت اقلام معیوب واقعی و غیر واقعی (جعلی) باشند. در این صورت داریم $\tilde{p} = \tilde{P}(E_1)$ و $\tilde{p}_e = \tilde{P}(A)$. مطابق با قانون جمع احتمال کل، روش باکلی [۵] و در نظر گرفتن $e_1 = P(A|E_2)$ و $e_2 = P(A^c|E_1)$ داریم:

$$\begin{aligned}\tilde{p}_e &= \tilde{p}(1 - e_2) + \tilde{q}e_1, \\ \tilde{p}_e[\lambda] &= \{p(1 - e_2) + qe_1 : s\}.\end{aligned}\quad (۱.۳)$$

با استفاده از رابطه (۳.۲)، λ - برش احتمال پذیرش فازی انباشته k ام در حضور خطاهای بازرسی $(\tilde{P}_{(a:k)(e)})$ برابر است با:

$$\tilde{P}_{(a:k)(e)}(\tilde{p}_e)[\lambda] = \left(\sum_{d=0}^{c_1} \tilde{P}_{d(e)}(\tilde{p}_e) \right)[\lambda] + \left(\sum_{d=c_1+1}^{c_2} \tilde{P}_{d(e)}(\tilde{p}_e) \right)[\lambda] \times \tilde{P}_{(a:k)(e)}^m(\tilde{p}_e)[\lambda], \quad (۲.۳)$$

که در آن $\tilde{P}_{d(e)}(\tilde{p}_e)$ ، احتمال فازی اینکه دقیقا d کالای معیوب در یک نمونه به اندازه n ، زمانی که نسبت اقلام معیوب غیر واقعی $\tilde{p}_e$ باشد، است.

۴. طرح FMDS(0, 1, 2) در حضور خطای بازرسی

منحنی مشخصه عملکرد^۳ (OC) طرح نمونه‌گیری ساده^۴ (SSP) با عدد پذیرش برابر صفر، در حالت کلاسیک و به ویژه در نزدیکی مبدأ، جایی که نسبت اقلام معیوب انباشته کم است، تقعر رو به بالا داشته و در چنین طرح‌هایی تنها با مشاهده یک قلم کالای معیوب، انباشته تحت بررسی، رد شده که به ضرر تولید کننده تمام می‌شود. طرح MDS با مقادیر $c_1 = 0$ و $c_2 = 1$ و $m = 2$ که به اختصار با نماد MDS(0,1,2) نشان داده می‌شود ضمن دارا بودن حجم نمونه

کم، موجب کم‌رنگ شدن ضعف‌های اشاره شده در طرح SSP می‌شود [۱۷]. از این رو در این بخش طرح $MDS(0,1,2)$ در محیط فازی که با نماد $FMDs(0,1,2)$ نشان داده می‌شود، مورد مطالعه قرار می‌گیرد.

همان‌طور که می‌دانیم یکی از ابزارهای مهم جهت مقایسه عملکرد طرح‌های مختلف، منحنی OC آن طرح‌ها می‌باشد. منحنی OC، تغییرات احتمال پذیرش را در برابر نسبت اقلام معیوب مختلف نشان می‌دهد [۱۵].

فرض کنید نسبت اقلام معیوب واقعی در انباشته یک عدد فازی مثلثی $\tilde{p} = (a_1, a_2, a_3)$ باشد. به منظور ترسیم منحنی OC در حضور خطاهای بازرسی e_1 و e_2 ، $\tilde{p}$ را به صورت $\tilde{p}_z$ بازنویسی می‌کنیم:

$$\begin{aligned}\tilde{p}_z &= (z, b_2 + z, b_3 + z), \\ \tilde{p}_z[\lambda] &= [z + b_2\lambda, b_3 + z - (b_3 - b_2)\lambda],\end{aligned}\quad (1.4)$$

که در آن $z \in [0, 1 - b_3]$ و $(i = 2, 3)$ ، $b_i = a_i - a_1$ است. به طور مشابه اگر $\tilde{p}_e$ به صورت $\tilde{p}_e = \tilde{p}_z(1 - e_2) + \tilde{q}_ze_1$ بازنویسی شود آنگاه

$$\tilde{p}_{ze}[\lambda] = [P_{ze}^L, P_{ze}^U], \quad (2.4)$$

که در آن

$$P_{ze}^L = (1 - e_2)(z + b_2\lambda) + e_1[(1 - z - b_3) + (b_3 - b_2)\lambda], \quad (3.4)$$

$$P_{ze}^U = (1 - e_2)[(z + b_3) + (b_2 - b_3)\lambda] + e_1[1 - z - b_2\lambda]. \quad (4.4)$$

فرض کنید $\tilde{P}_{(a:k)(ze)}[\lambda]$ ، λ -برش احتمال پذیرش فازی انباشته k ام در حضور خطاهای بازرسی با نسبت اقلام معیوب برابر $\tilde{p}_{ze}$ باشد. در آن صورت با توجه به رابطه (۲.۳) داریم:

$$\tilde{P}_{(a:k)(ze)}(\tilde{p}_{ze})[\lambda] = \tilde{P}_{0(ze)}(\tilde{p}_{ze})[\lambda] + \tilde{P}_{1(ze)}(\tilde{p}_{ze})[\lambda] \times \tilde{P}_{(a:k)(ze)}^2(\tilde{p}_{ze})[\lambda], \quad (5.4)$$

که در آن

$$\tilde{P}_{0(ze)}(\tilde{p}_{ze})[\lambda] = \{(1 - p_{ze})^n : s\} = [P_{0(ze)}^L[\lambda], P_{0(ze)}^U[\lambda]],$$

$$\tilde{P}_{1(ze)}(\tilde{p}_{ze})[\lambda] = \{np_{ze}(1 - p_{ze})^{n-1} : s\} = [P_{1(ze)}^L[\lambda], P_{1(ze)}^U[\lambda]],$$

$$P_{0(ze)}^L[\lambda] = \min\{(1 - p_{ze})^n : s\},$$

$$P_{0(ze)}^U[\lambda] = \max\{(1 - p_{ze})^n : s\},$$

$$P_{1(ze)}^L[\lambda] = \min\{np_{ze}(1 - p_{ze})^{n-1} : s\},$$

$$P_{1(ze)}^U[\lambda] = \max\{np_{ze}(1 - p_{ze})^{n-1} : s\},$$

و

$$\begin{aligned}\tilde{P}_{(a:k)(ze)}(\tilde{p}_{ze})[\lambda] &= \left\{ \frac{1 - \sqrt{1 - 4np_{ze}(1 - p_{ze})^{2n-1}}}{2np_{ze}(1 - p_{ze})^{n-1}} : s \right\} \\ &= [P_{(a:k)(ze)}^L[\lambda], P_{(a:k)(ze)}^U[\lambda]],\end{aligned}\quad (6.4)$$

به طوریکه $P_{(a:k)(ze)}^L[\lambda]$ و $P_{(a:k)(ze)}^U[\lambda]$ به ترتیب می‌نیم و ماکزیمم $\tilde{P}_{(a:k)(ze)}(\tilde{p}_{ze})[\lambda]$ می‌باشد.

مثال ۱.۴. فرض کنید در یک فرایند تولیدی، نسبت اقلام معیوب واقعی برابر عدد فازی $\tilde{p}$ بوده و بازرسی در حضور خطاهای بازرسی $e_1 = 0.01$ و $e_2 = 0.15$ انجام پذیرد. مطابق معاهده نوشته شده بین مشتری و تولیدکننده انباشته k ام در صورتی پذیرفته می‌شود که در یک نمونه ۱۰ تایی گرفته شده از آن، هیچ کالای معیوبی مشاهده نشود و یا در صورت مشاهده یک قلم کالای معیوب، هر دو انباشته متوالی بعدی فاقد اقلام معیوب باشند. در آن صورت مطابق روابط (۱.۴) و (۲.۴) داریم:

$$\tilde{p}_{ze}[0] = [0.0098 + 0.84z, 0.027 + 0.84z]$$

$$\tilde{p}_z[0] = [z, z + 0.02], z \in [0, 0.98].$$

با استفاده از رابطه (۶.۴)، برخی از مقادیر $\tilde{P}_{(a:k)z}[0]$ و $\tilde{P}_{(a:k)(ze)}[0]$ در جدول ۱ گزارش شده است. مطابق جدول ۱ ملاحظه می‌شود که کران‌های بازه ۰-برش احتمالات پذیرش فازی در حضور و عدم حضور خطاهای بازرسی کاهش می‌یابند هرگاه نسبت اقلام واقعی و غیر واقعی افزایش می‌یابند.

جدول ۱: مقادیر ۰-برش احتمالات پذیرش انباشته در حضور و عدم حضور خطاهای بازرسی در طرح FMDS(0,1,2).

z	$\tilde{p}_z[0]$	$\tilde{p}_{ze}[0]$	$\tilde{P}_{(a:k)z}[0]$	$\tilde{P}_{(a:k)(ze)}[0]$
0	[0, 0.02]	[0.0098, 0.0270]	[0.8307, 0.9995]	[0.7649, 0.9192]
0.01	[0.01, 0.03]	[0.0182, 0.0354]	[0.7358, 0.9174]	[0.6815, 0.8461]
0.02	[0.02, 0.04]	[0.0266, 0.0438]	[0.6345, 0.8296]	[0.5951, 0.7679]
0.03	[0.03, 0.05]	[0.0350, 0.0522]	[0.5318, 0.7346]	[0.5097, 0.6847]
0.04	[0.04, 0.06]	[0.0434, 0.0606]	[0.4356, 0.6322]	[0.4300, 0.5983]
0.05	[0.05, 0.07]	[0.0518, 0.0690]	[0.3515, 0.5306]	[0.3591, 0.5128]

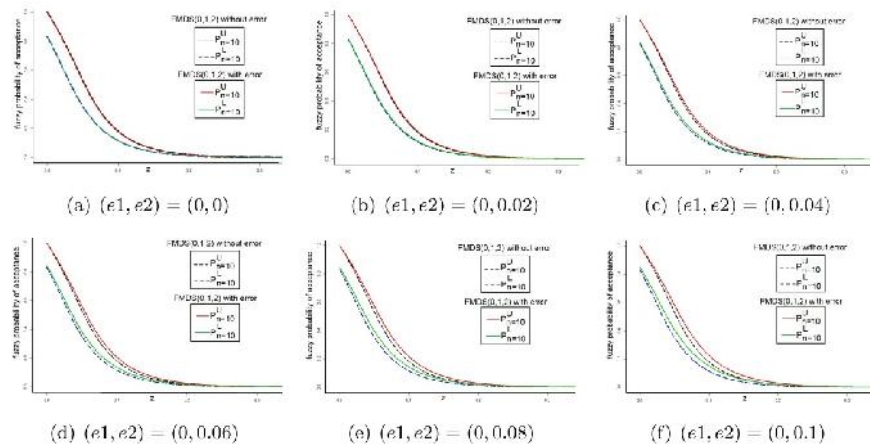
در ادامه به بررسی تاثیر مقادیر مختلف خطاهای بازرسی بر روی منحنی OC طرح پیشنهادی در صورتی که مقدار $\tilde{p}$ ثابت باشد، می‌پردازیم.

فرض کنید $\tilde{p}_e$ و $\tilde{p}_e^*$ به ترتیب نسبت اقلام معیوب انباشته در حضور خطاهای بازرسی (e_1, e_2) و (e_1^*, e_2^*) باشند.

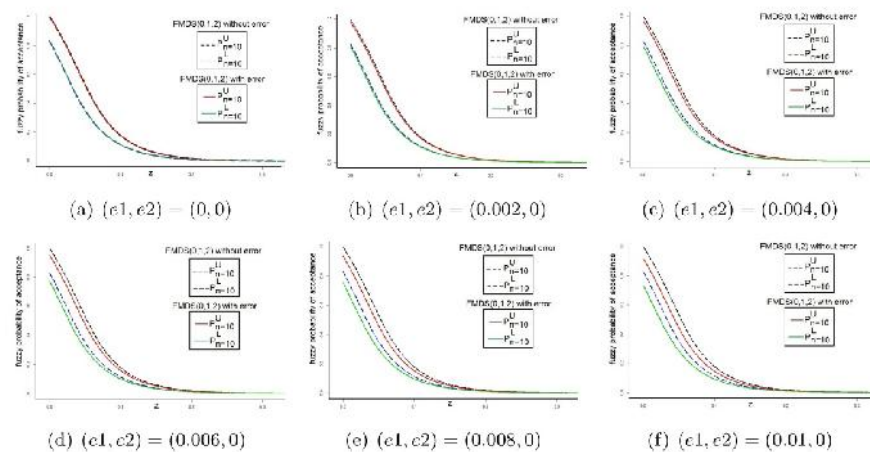
با تغییر خطاهای بازرسی از (e_1, e_2) به (e_1^*, e_2^*) ، امکان وقوع یکی از دو حالت الف) $\tilde{p}_e \leq \tilde{p}_e^*$ یا ب) $\tilde{p}_e > \tilde{p}_e^*$ وجود دارد. به وضوح اگر $\tilde{p}(e_2^* - e_2) \leq (1 - \tilde{p})(e_1^* - e_1)$ باشد آنگاه حالت الف) رخ داده و در غیر این صورت با حالت ب) مواجه خواهیم بود.

با ثابت گرفتن مقدار e_1 و افزایش e_2 ، مقدار نسبت اقلام معیوب به اشتباه، مقداری کمتر از واقعیت گزارش می‌شود (حالت ب)) در نتیجه، احتمال پذیرش فازی انباشته، افزایش می‌یابد. حال اگر، e_2 ثابت و مقدار e_1 افزایش یابد (حالت الف))، به تبع آن احتمال پذیرش فازی انباشته کاهش می‌یابد.

شکل ۱، ۰-برش منحنی OC طرح FMDS(0,1,2) را به ازای $e_1 = 0$ و $e_2 = 0, 0.02, 0.04, 0.06, 0.08, 0.1$ نشان می‌دهد. ۰-برش منحنی OC طرح پیشنهادی وقتی $e_1 = 0, 0.002, 0.004, 0.006, 0.008, 0.01$ و $e_2 = 0$ مطابق شکل ۲ ترسیم شده است.



شکل ۱: ۰-برش منحنی OC طرح FMDS(0,1,2) وقتی $e_1 = 0$ و e_2 تغییر می‌کند.



شکل ۲: ۰-برش منحنی OC طرح FMDS(0,1,2) وقتی $e_2 = 0$ و e_1 تغییر می‌کند.

در هر دو شکل ۱ و ۲، منحنی‌های توپر و نقطه‌چین به ترتیب ۰-برش منحنی OC طرح پیشنهادی در حضور و بدون حضور خطاهای بازرسی با اندازه نمونه $n = 10$ است. از شکل ۱ مشاهده می‌شود که با ثابت گرفتن مقدار e_1 و افزایش e_2 میزان احتمال پذیرش فازی انباشته افزایش می‌یابد. زیرا در این حالت مقدار نسبت اقلام معیوب انباشته به اشتباه

کاهش یافته است. این مسأله اگر چه از دیدگاه تولیدکننده مطلوب است اما از جهت مشتری به ویژه به ازای مقادیر بزرگتر $\bar{p}$ ، اتفاق خوشایندی نمی باشد. از سوی دیگر، شکل ۲ نشان می دهد که مقدار احتمال پذیرش فازی، با افزایش e_1 ، کاهش می یابد. زیرا در چنین حالتی مقدار نسبت اقلام معیوب به اشتباه، زیاد گزارش شده است که از دیدگاه تولید کننده، پدیده مطلوبی نمی نماید.

۵. نتیجه

در مطالعه حاضر، ضمن معرفی طرح MDS در محیط فازی، به توسعه طرح پیشنهادی در حضور خطاهای بازرسی پرداخته شد. به منظور بررسی تاثیر خطاهای بازرسی بر روی طرح پیشنهادی، نظریه مجموعه های فازی بکار گرفته شد. نتایج به دست آمده نشان دادند که طرح پیشنهادی دربرگیرنده طرح MDS موجود در حالت کلاسیک بوده و قدرت عملکرد آن در تشخیص کیفیت خوب و بد انباشته ای از تولیدات، تحت تاثیر خطاهای بازرسی قرار می گیرد که از نظر تولید کننده و مصرف کننده مطلوب نمی باشد.

مراجع

1. Baloui Jamkhaneh, E., Sadeghpour Gildeh, B. and Yari, G. (2011a). Acceptance single Sampling plan with fuzzy parameter, Iranian Journal of Fuzzy Systems, 8(2), 47-55.
2. Baloui Jamkhaneh, E., Sadeghpour Gildeh, B. and Yari, G. (2011b). Inspection error and its effects on single sampling plans with fuzzy parameters, Structural and multidisciplinary Optimization, 43(4), 555-560.
3. Baloui Jamkhaneh, E. and Sadeghpour Gildeh, B. (2012). Acceptance double sampling plan using fuzzy poisson distribution, World Applied Sciences Journal, 16(11), 1578-1588.
4. Baloui Jamkhaneh, E. and Sadeghpour Gildeh, B. (2013). Sequential sampling plan using fuzzy SPRT, Journal of Intelligent and Fuzzy Systems, 25(3), 785-791.
5. Buckley, J. J. (2006). Fuzzy probability and statistics, Berlin Heidelberg: Springer Verlag.
6. Case, K. E., Benett, G. K. and Schmidt, J. W. (1975). The effect on inspection error on average outgoing quality, Journal of Quality Technology, 7(1), 28-33.
7. Chakraborty, T. K. (1992). A class of single sampling plans based on fuzzy optimization, Quality Control and Applied Statistics, 37(7), 359-362.
8. Chen, C. H. and Chou, C. Y. (2003). Economic specification limits under the inspection error, Journal of the Chinese Institute of Industrial Engineers, 20(1), 9-13.
9. Dorris, A. L. and Foote, B. L. (1978). Inspection errors and statistical quality control, AIIE Transactions, 10(2), 184-192.
10. Duffuaa S. O. and El-Ga'aly A., "Impact of inspection errors on the formulation of a multi-objective optimization process targeting model under inspection sampling plan", Computers and Industrial Engineering, Vol. 80, PP. 254-260, 2015.
11. Fallahnezhad M. S. and Yousefi Babadi A., "A new acceptance sampling plan using bayesian approach in the presence of inspection errors", Transactions of the Institute of Measurement and Control, Vol. 37, No. 9, PP. 1060-1073, 2015.

12. Grzegorzewski, P. (2002). Acceptance sampling plans by variables for vague data, *Soft Methods in Probability, Statistics and Data Analysis*. Physica-Verlag HD.
13. Kanagawa, A. and Ohta, H. (1990). A design for single sampling attribute plan based on fuzzy sets theory, *Fuzzy Sets and Systems*, 37(2), 173-181.
14. Latha, M. and Subbiah, K.: Selection of bayesian multiple deferred state (BMDS-1) sampling plan based on quality regions. *International Journal of Recent Scientific Research*. 6(5), 3864-3867 (2015).
15. Montgomery, D. C. (1991). *Introduction to statistical quality control*, Wiley, New York.
16. Osanaiye, P. A. and Alebiosu S. A. (1988). Effects of industrial inspection errors on some plans that utilise the surrounding lot information, *Journal of Applied Statistics*, 15(3), 295-304.
17. Soundararajan, V. and Vijayaraghavan, R. (1990). Construction and selection of multiple dependent (deferred) state sampling plan, *Journal of Applied Statistics*, 17(3), 397-409.
18. Subramain, K. and Haridoss, V. (2012). Development of multiple deferred state sampling plan based on minimum risks using the weighted poisson distribution for given acceptance quality level and limiting quality level, *International Journal of Quality Engineering and Technology*, 3(2), 168-180.
19. Tamaki, F., Kanagawa, A. and Ohta, H. (1991). A Fuzzy Design of Sampling Inspection Plans by Attributes, *Japanese Journal of Fuzzy Theory and Systems*, 3(4), 315-327.
20. Tong, X. and Wang, Z. (2012). Fuzzy acceptance sampling plans for inspection of geospatial data with ambiguity in quality characteristics, *Computers and Geosciences*, 48(11), 256-266.
21. Turanoglu, E., Kaya, I. and Kahraman, C. (2012). Fuzzy acceptance sampling and characteristic curves, *International Journal of Computational Intelligence Systems*, 5(1), 13-29.
22. Venkateh, A. and Elango, S. (2014). Acceptance sampling for the influence of TRH using crisp and fuzzy gamma distribution, *Aryabhatta Journal of Mathematics and Informatics*, 6(1), 119-124.
23. Wortham, A. W. and Baker, R. C. (1976). Multiple deferred state sampling inspection, *International Journal of Production Research*, 14(6), 719-731.
24. Wu, C-W., Lee, A. and Chen, Y.: A novel lot sentencing method by variables inspection considering multiple dependent state. *Quality and Reliability Engineering International*. 32 (3), 985-994 (2016).
25. Yan, A., Liu, S. and Dong, X.: Designing a multiple dependent state sampling plan based on the coefficient of variation. *SpringerPlus*. 5(1), 1447 (2016).



کیفیت فازی: نمودارهای کنترل شوهارت

عباس پرچمی *

دانشگاه شهید باهنر کرمان، دانشکده ریاضی و کامپیوتر، بخش آمار
parchami@uk.ac.ir

چکیده. پس از معرفی کیفیت فازی در این مقاله به تعمیم نمودارهای کنترل شوهارت مبتنی بر کیفیت فازی پرداخته شده است. این نمودارها شامل نمودار کنترل فازی تعداد اقلام معیوب، نمودار کنترل فازی نسبت اقلام معیوب و نمودار کنترل فازی تعداد نقصها در واحد بازرسی می باشند. وجه مشترک تمامی این نمودارها، استفاده از آماره‌ی تعداد اقلام معیوب/ناقص است که تابعی از یک مشخصه کیفی می باشد و البته این مشخصه کیفی، مبتنی بر کیفیت فازی بصورت کمی اندازه گیری می شود. وقتی استفاده از این نمودارهای کنترل فازی توجیه عملی دارد که درجه‌ی بی کیفیت (میزان عیب) تمامی اقلام ناقص با یکدیگر مساوی و یکسان نباشند.

۱. پیشگفتار

حدود کنترل بالا و پایین (حدود عمل) که به ترتیب با LCL و UCL نشان داده می شوند به عنوان تابعی از میزان تغییرپذیری طبیعی فرآیند، محدوده‌ای را تشکیل می دهند که در صورت تحت کنترل بودن فرآیند تولیدی، مقدار آماره‌ی مربوط به کیفیت کالاهای تولیدی، در داخل آن محدوده به صورت کاملاً تصادفی نوسان خواهد داشت. اگر یک مشخصه کیفیت به وسیله‌ی آماره‌ی W اندازه گیری شود و همچنین میانگین و انحراف معیار W به ترتیب μ_W و σ_W باشند، آن گاه مدل کلی نمودار کنترل شوهارت به صورت زیر خواهد بود:

$$\begin{aligned} LCL &= \mu_W - k \sigma_W \\ CL &= \mu_W \\ UCL &= \mu_W + k \sigma_W \end{aligned} \quad (1.1)$$

که در آن، فاصله‌ی حدود کنترل LCL و UCL از خط مرکزی CL ، k برابر انحراف معیار W در نظر گرفته شده است [۴، ۱]. پیش نیاز برقراری هر سیستم کنترل، فراهم کردن اطلاعات از طریق اندازه گیری متغیر تحت بررسی است. لذا بعد از تعیین متغیر تصادفی مورد نظر که مشخصه کیفیت را با آن نمایش می دهیم، باید یکی از دو نوع مقیاس زیر به منظور اندازه گیری و ثبت داده‌های مربوط به مشخصه کیفیت مورد نظر انتخاب گردد:

الف) مقیاس کیفی: کیفیت کالا در مقیاس کیفی به صورت نظری و در قالب مرغوب/نامرغوب

واژگان کلیدی. کیفیت فازی، نمودار کنترل شوهارت، مقیاس کیفی.
* سخنران

ارزیابی و دسته‌بندی می‌شود. مثلاً رنگ کالا مطابق استاندارد نبوده و کالا نامرغوب است. **ب) مقیاس کمی:** کیفیت در مقیاس کمی به صورت پیوسته اندازه‌گیری می‌شود. مثلاً طول، جرم و دمای کالای تولیدی، هر کدام به عنوان یک مشخصه کیفیت، می‌توانند به وسیله‌ی یک مقیاس کمی اندازه‌گیری شوند.

می‌دانیم که دامنه‌ی مربوط به یک مشخصه‌ی کمی معمولاً بازه‌ای پیوسته در نظر گرفته می‌شود و این متغیر با توجه به ماهیتش می‌تواند بطور پیوسته در این محدوده تغییر کند. لذا مقادیر ثبت شده از یک متغیر کمی حاوی اطلاعات زیاد و دقیقی از متغیر مربوطه می‌باشند. اما در مقابل، مقادیر مشاهده شده از یک متغیر کیفی که به کمک یک مقیاس کیفی (وصفی) اندازه‌گیری شده باشند، اطلاعات بسیار کمی را از متغیر مربوطه در بر دارند؛ زیرا این متغیرها تنها مجاز به اختیار کردن دو مقدار صفر و یک (که بترتیب به منزله‌ی خوب/بد، با کیفیت/بی‌کیفیت، قابل قبول/غیرقابل قبول، و یا سالم/معیوب هستند) می‌باشند. در این مقاله سعی شده است تا این کمبود اطلاعات نهفته در داده‌های متغیر کیفی، به کمک تعمیمی بر نمودارهای کیفی کنترل که بر اساس نظریه‌ی مجموعه‌های فازی مطرح شده است، برطرف گردد. بدین منظور، در بخش ۲ قصد داریم تا نگاهی به مفهوم کیفیت فازی برای استفاده در نمودارهای کنترل شوهارت بیان‌داریم. در بخش ۳، نمودارهای کنترل کیفیت فازی را مبتنی بر کیفیت فازی تعمیم می‌دهیم و سپس حدود کنترل آنها را در بخش ۴ بر اساس داده‌ها برآورد می‌کنیم.

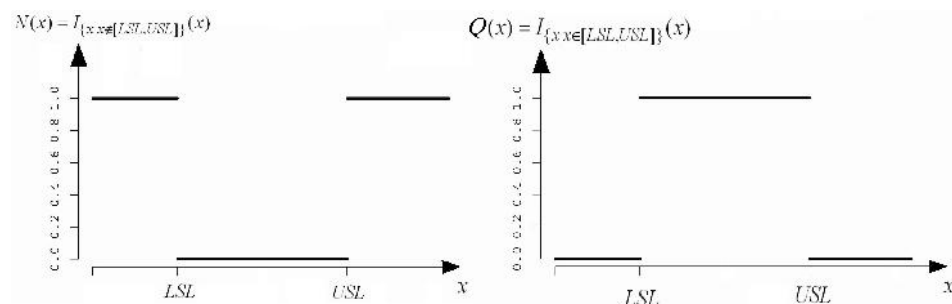
۲. کیفیت فازی

اگر مقدار مشخصه کیفیت اندازه‌گیری شده‌ی یک قلم محصول تولیدی با استانداردها و حدود مشخصه فنی تطابق داشته باشد و بین حدود مشخصه فنی بالا و پایین^۱ که به ترتیب با USL و LSL نشان داده می‌شوند قرار گیرد، آنگاه آن محصول سالم و در غیر این صورت محصول معیوب شناخته می‌شود. از این رو می‌توان درجه کیفیت و درجه بی‌کیفیتی مقدار مشخصه کیفیت x را به ترتیب با توابع $Q(x)$ و $N(x) = 1 - Q(x)$ نشان داد که در آنها

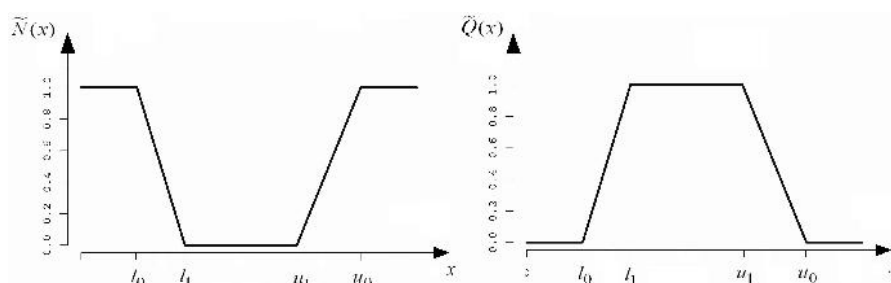
$$Q(x) = \begin{cases} 1 & \text{اگر } x \in [LSL, USL] \\ 0 & \text{اگر } x \notin [LSL, USL] \end{cases} \quad (۱.۲)$$

تابع نشانگر بازه‌ی $[LSL, USL]$ است (شکل ۱ را ملاحظه کنید). مشخصه‌ی کیفیت X با مجموعه‌ی مرجع (دامنه‌ی) χ را در نظر بگیرید. همان طور که مطرح شد، کیفیت معمولی (غیرفازی) بوسیله‌ی بازه‌ی $[LSL, USL]$ مشخص می‌شود که می‌توان آن را با تابع نشانگر $I_{[LSL, USL]}(x)$ نیز نمایش داد. برای یک کالای خاص با حدود مشخصه معلوم، $x \in [LSL, USL]$ (یا به طور معادل $I_{[LSL, USL]}(x) = 1$)، به معنای باکیفیت بودن آن کالا است و در مقابل $I_{[LSL, USL]}(x) = 0$ به معنای بی‌کیفیتی آن کالا است. همانند هر مجموعه‌ی فازی که بوسیله‌ی تابع عضویتش تعریف و مشخص می‌شود، کیفیت فازی نیز بوسیله‌ی تابع عضویتش، که تعمیمی از تابع نشانگر $I_{[LSL, USL]}(x)$ است، تعریف می‌شود [۷].

^۱Upper and Lower Specification Limits



شکل ۱: توابع نشانگر مجموعه‌ی کالاهای باکیفیت (تصویر سمت راست) و مجموعه‌ی کالاهای بی‌کیفیت (تصویر سمت چپ) در کیفیت معمولی (غیر فازی).



شکل ۲: توابع عضویت مجموعه‌ی فازی کالاهای باکیفیت (تصویر سمت راست) و مجموعه‌ی فازی کالاهای بی‌کیفیت (تصویر سمت چپ) در کیفیت فازی.

تعریف ۱.۲. تابع عضویت کیفیت فازی $\tilde{Q}$ تابعی از χ به بازه‌ی $[0, 1]$ است به‌طوری‌که اگر x مقدار مشخصه کیفیت مشاهده‌شده برای یک کالا باشد آن‌گاه $\tilde{Q}(x)$ میزان کیفیت آن کالا را نشان دهد. به عبارت دیگر، $\tilde{Q} \in F(\chi)$ کیفیت فازی برای مشخصه کیفیت X است اگر و تنها اگر $\tilde{Q}(x)$ نشان دهنده‌ی میزان کیفیت آن کالا به ازای هر $x \in \chi$ باشد که در آن $F(\chi)$ مجموعه تمام مجموعه‌های فازی روی χ است.

بنابراین در کیفیت فازی، مقدار $\tilde{Q}(x)$ میزان کیفیت یک قلم محصولی می‌باشد که مشخصه کیفیت آن برابر x است و مقدار $\tilde{Q}(x) = 1 - \tilde{N}(x)$ میزان بی‌کیفیتی آن محصول می‌باشد (شکل ۲ را ملاحظه کنید).

در کیفیت فازی حالت‌های بیشتری نسبت به دو حالت محدود صفر و یک (سالم و معیوب) می‌توان برای کیفیت یک محصول در نظر گرفت و براساس نمودارهای کنترل قضاوت منصفانه‌تری در مورد تحت کنترل بودن یک فرایند تولیدی ارائه نمود. در واقع اگر ملاک کیفیت فازی را در رسم نمودار کنترل یک کالا لحاظ کنیم، آن‌گاه داده‌های ثبت شده برای مشخصه‌ی کیفیت (که در حقیقت درجات عضویت به مجموعه کیفیت فازی $\tilde{Q}$ می‌باشند) حاوی اطلاعات بیشتری،

^۲Fuzzy Quality

نسبت به کیفیت غیرفازی هستند. زیرا در این حالت، دیگر داده‌های ثبت شده لزوماً مقادیر صفر یا یک نبوده و مجازند تا هر عددی از بازه‌ی $[0, 1]$ را به عنوان درجه‌ی کیفیت کالا اختیار کنند و همین موضوع سبب ارتقاء ارزش داده‌های ثبت شده برای یک متغیر کیفی در حالت کیفیت فازی می‌گردد. بدیهی است که استنتاج به کمک این داده‌ها و بر اساس نمودارهای کنترل شوهارت، جهت بررسی تحت کنترل بودن و یا تحت کنترل نبودن فرایند، بسیار معتبرتر از روش‌های سنتی و متداول خواهد بود. در بخش بعد روش رسم این گونه نمودارهای کنترل و محاسبه‌ی حدود کنترل برای این نمودارها را مورد بحث قرار می‌دهیم.

۳. نمودارهای کنترل کیفیت فازی

برخی از نمودارهای کنترل شوهارت مبتنی بر آماره‌ای ترسیم می‌شوند که مفهوم «تعداد» را در بر دارد. بعنوان نمونه می‌توان به نمودار کنترل تعداد اقلام معیوب (نمودار کنترل np) و همچنین نمودار کنترل تعداد نقص‌ها در واحد بازرسی (نمودار کنترل C) اشاره کرد. در کیفیت معمولی و غیر فازی، مفهوم «تعداد» را می‌توان به کمک مفهوم تابع نشانگر بصورت زیر بازنویسی نمود. فرض کنید $X_1, X_2, \dots, X_n$ یک نمونه‌ی تصادفی n تایی از مشخصه‌ی کیفی X باشد، آنگاه با در نظر گرفتن مفهوم کیفیت غیرفازی که در بخش ۲ معرفی شد، $N(X_1), \dots, N(X_n)$ نیز یک نمونه تصادفی n تایی از متغیر تصادفی $N(X)$ است بطوریکه به ازای $i = 1, \dots, n$ داریم $N(X_i) \in \{0, 1\}$. در اینصورت می‌توان تعداد بی‌کیفیتی‌ها (مثلاً تعداد نقص‌ها، یا تعداد اقلام معیوب) در یک نمونه‌ی تصادفی n تایی را بصورت مجموع n عدد صفر یا یک (که بترتیب نشانگر باکیفیتی و بی‌کیفیتی کالاها می‌باشند)، یعنی $\sum_{i=1}^n N(X_i)$ ، نشان داد. بر اساس این ترفند که معمولاً برای تعمیم تعاریف معمولی به تعاریف فازی متداول و شایع است، قضیه زیر خط مرکزی و حدود کنترل نمودارهای کنترل شوهارت را مبتنی بر کیفیت فازی تعمیم می‌دهد. این تعمیم برای حالتی است که آماره‌ی نمودارهای کنترل بر اساس مشخصه کیفیتی با مفهوم «تعداد» باشد.

قضیه ۱.۳. فرض کنید مشخصه‌ی کیفی X به وسیله‌ی آماره‌ی W ، که به نوعی مفهوم تعداد را در بر دارد، اندازه‌گیری می‌شود. در این صورت، مدل کلی حدود کنترل و خط مرکزی در نمودارهای کنترل شوهارت مبتنی بر کیفیت فازی $\tilde{Q}$ به صورت زیر می‌باشد:

$$\begin{aligned} LCL &= n \left(1 - \mu_{\tilde{Q}} - k \frac{\sigma_{\tilde{Q}}}{\sqrt{n}} \right) \\ CL &= n (1 - \mu_{\tilde{Q}}) \\ UCL &= n \left(1 - \mu_{\tilde{Q}} + k \frac{\sigma_{\tilde{Q}}}{\sqrt{n}} \right) \end{aligned} \quad (۱.۳)$$

که در آنها

$$\begin{aligned} \mu_{\tilde{Q}} &= \int_{-\infty}^{+\infty} \tilde{Q}(x) f(x) dx \\ \sigma_{\tilde{Q}} &= \left\{ \int_{-\infty}^{+\infty} [\tilde{Q}(x) - \mu_{\tilde{Q}}]^2 f(x) dx \right\}^{\frac{1}{2}} \end{aligned} \quad (۲.۳)$$

و f تابع چگالی احتمال مشخصه کیفیت X است. در صورتی که مشخصه کیفی X گسسته باشد، f تبدیل به تابع جرم احتمال و در نتیجه انتگرال‌های فوق به مجموع‌یایی تبدیل می‌شوند.

نکته ۲.۳. در قضیه ۱.۳، حدود کنترل نمودارهای فازی به نحوی معرفی شده‌اند که آماره‌ی آنها مبتنی بر یک مشخصه کیفی بوده و به مفهوم «تعداد» نیز می‌باشد؛ مانند نمودار کنترل فازی تعداد اقلام معیوب، نمودار کنترل فازی نسبت اقلام معیوب و نمودار کنترل فازی تعداد نقص‌ها در واحد بازرسی، که از این پس جهت سهولت ما آنها را بترتیب با نمودارهای کنترل $\tilde{np}$ ، $\tilde{p}$ و $\tilde{C}$ نشان می‌دهیم.

نکته ۳.۳. اگر کیفیت به صورت معمولی و غیرفازی در نظر گرفته شود، یعنی تابع عضویت مجموعه‌ی فازی کالاهای باکیفیت تابع نشانگر Q در نظر گرفته شود، آنگاه بوسیله این تابع نشانگر محصولات در دو گروه سالم و معیوب ارزشیابی و دسته‌بندی شده و مجموع درجات بی‌کیفیتی‌ها یعنی آماره‌ی $\tilde{W} = \sum_{j=1}^n \tilde{N}(X_j) = \sum_{j=1}^n (1 - \tilde{Q}(X_j))$ تبدیل به آماره‌ی W یعنی تعداد بی‌کیفیتی‌ها (تعداد نقص‌ها/معیوب‌ها) در فرایندهای غیرفازی می‌شود. از آنجا که اساس معرفی حدود کنترل در قضیه ۱.۳ جایگذاری آماره‌ی $\tilde{W}$ بجای آماره‌ی W است، لذا می‌توان ادعا کرد که در کیفیت غیرفازی، نمودارهای کنترل معرفی شده‌ی $\tilde{p}$ ، $\tilde{np}$ و $\tilde{C}$ ، بترتیب تبدیل به نمودارهای کنترل معمولی و متداول p ، np و C می‌شوند [۱].

۴. برآورد حدود کنترل و خط مرکزی مبتنی بر کیفیت فازی

در این بخش قصد داریم تا بر اساس قضیه ۱.۳، روند برآورد حدود کنترل و همچنین روش رسم نمودارهای کنترل فازی $\tilde{p}$ ، $\tilde{np}$ و $\tilde{C}$ را به طور دقیق‌تر و مبتنی بر کیفیت فازی $\tilde{Q}$ بررسی کنیم.

۱.۴. نمودار کنترل فازی تعداد اقلام معیوب (نمودار کنترل $\tilde{np}$). تعداد اقلام معیوب (بی‌کیفیت) را از یک زمان به زمان دیگر، یا از یک نمونه به نمونه‌ی دیگر، می‌توان بر روی نمودار کنترل $\tilde{np}$ و بر اساس کیفیت فازی رسم کرد. تنها وقتی استفاده از نمودار کنترل $\tilde{np}$ بجای نمودار کنترل np در عمل توجیه دارد که درجه‌ی بی‌کیفیتی (میزان عیب) تمامی اقلام معیوب از دیدگاه کاربر مساوی و یکسان نباشد. برای رسم نمودار کنترل $\tilde{np}$ بر اساس اطلاعات m نمونه‌ی n تایی، کافیت نقاط $(i, \tilde{w}_i)$ را به ازای $i = 1, 2, \dots, m$ بر روی یک نمودار دوبعدی رسم کرده و سپس این نقاط را متوالیاً به یکدیگر متصل سازیم بطوریکه

$$\tilde{w}_i = \sum_{j=1}^n \tilde{N}(x_{ij}) = \sum_{j=1}^n (1 - \tilde{Q}(x_{ij})) \quad (1.4)$$

مجموع درجات بی‌کیفیتی در نمونه‌ی i ام، مقدار مشخصه کیفیت مربوط به i امین عضو نمونه‌ی i ام، و $\tilde{Q}$ تابع عضویت کیفیت فازی است. به این ترتیب، اگر تمامی نقاط رسم شده بین حدود کنترل قرار داشته و هیچ گونه رفتار سیستماتیکی از خود نشان ندهند، آنگاه می‌توان

تحت کنترل بودن آماری مربوطه را نتیجه گرفت و در غیر اینصورت فرایند از کنترل خارج بوده و نیازمند اقدامات اصلاحی می باشد.

نکته‌ی قابل توجه آن است که x در فرمول‌های (۲.۳) تعداد اقلام معیوب نیست بلکه مشخصه کیفیت (مثلاً طول خَش یا مساحت لکه بر بدنه‌ی کالای تولیدی) است. گر چه تابع چگالی احتمال این مشخصه کیفیت (یعنی $f(x)$) برای کاربران معمولاً مجهول است، اما برای ترسیم نمودار کنترل شوهارت و تشخیص حدود کنترل تنها کفایت مقادیر میانگین و انحراف معیار $\bar{Q}(X)$ که در قضیه ۱.۳ با $\mu_{\bar{Q}}$ و $\sigma_{\bar{Q}}$ نمادین شده‌اند را بر اساس داده‌ها برآورد کرد که این کار توسط روابط

$$\begin{aligned}\bar{\bar{Q}} &= \sum_{i=1}^m \bar{Q}_i / m \\ s_{\bar{Q}} &= \sqrt{\sum_{i=1}^m s_i^2 / m}\end{aligned}\quad (2.4)$$

به راحتی انجام پذیر است بطوریکه $\bar{Q}_i = \frac{\sum_{j=1}^n \bar{Q}(x_{ij})}{n}$ و $s_i = \sqrt{\frac{\sum_{j=1}^n [\bar{Q}(x_{ij}) - \bar{Q}_i]^2}{n-1}}$ به ترتیب برآوردهای میانگین و انحراف معیار i امین نمونه هستند. لذا با توجه به قضیه ۱.۳، برآورد حدود کنترل نمودار $\bar{n}p$ برابرند با

$$\begin{aligned}\widehat{LCL} &= n \left(1 - \bar{\bar{Q}} - k \frac{s_{\bar{Q}}}{\sqrt{n}} \right) \\ \widehat{CL} &= n \left(1 - \bar{\bar{Q}} \right) \\ \widehat{UCL} &= n \left(1 - \bar{\bar{Q}} + k \frac{s_{\bar{Q}}}{\sqrt{n}} \right)\end{aligned}\quad (3.4)$$

۲.۴. نمودار کنترل فازی نسبت اقلام معیوب (نمودار کنترل $\bar{p}$). بر اساس کیفیت فازی، نسبت و یا درصد اقلام معیوب را از یک نمونه به نمونه‌ی دیگر می توان بر روی نمودار کنترل فازی نسبت اقلام معیوب رسم کرد. برای رسم نمودار کنترل $\bar{p}$ بر اساس اطلاعات m نمونه‌ی n تایی، باید نقاط $(i, \frac{\bar{w}_i}{n})$ را به ازای $i = 1, 2, \dots, m$ بر روی نمودار رسم کرده و سپس این نقاط را متوالیاً به یکدیگر متصل کنیم بطوریکه $\bar{w}_i$ مجموع درجات بی کیفیتی در نمونه‌ی i ام است و مطابق رابطه‌ی (۱.۴) بدست می آید.

برای محاسبه‌ی حدود کنترل نمودار فازی نسبت اقلام معیوب (نمودار کنترل $\bar{p}$)، کافی است تا حدود نمودار کنترل $\bar{n}p$ را تقسیم بر حجم نمونه یعنی n کنیم تا حدود نمودار کنترل $\bar{p}$ بدست آیند. بنابراین در این حالت برآورد حدود کنترل و خط مرکزی نمودار $\bar{p}$ بصورت زیر می باشند که در آن‌ها $\bar{Q}$ و $s_{\bar{Q}}$ از روابط (۲.۴) بدست می آیند

$$\widehat{LCL} = 1 - \bar{\bar{Q}} - k \frac{s_{\bar{Q}}}{\sqrt{n}}$$

$$\begin{aligned}\widehat{CL} &= 1 - \overline{\overline{Q}} \\ \widehat{UCL} &= 1 - \overline{\overline{Q}} + k \frac{s_{\overline{Q}}}{\sqrt{n}}\end{aligned}\quad (4.4)$$

گفتنی است که در مراجع [۵، ۶] ایده‌ی مشابه قضیه ۱۰۳ اما فقط برای نمودار کنترل $\tilde{p}$ مطرح شده است.

۳.۴. نمودار کنترل فازی تعداد نقص‌ها در واحد بازرسی (نمودار کنترل $\tilde{C}$). گاهی ممکن است وجود نقص در محصول آن را از «سالم» بودن نیندازد، بلکه کیفیت آن را کاهش دهد، مثلاً وجود مقداری ناصافی یا لک در بدنه اتومبیل موجب رد شدن آن نمی‌شود ولی کیفیت آن را پایین می‌آورد. وجود زدگی در پارچه یا کاشی نمونه‌های دیگری از این نوع مشخصه می‌باشند. در چنین مواردی معمولاً از نمودارهای کنترل تعداد نقص‌ها استفاده می‌شود که در آن تمامی نقص‌ها هم‌ارزش و یکسان در نظر گرفته شده و تحت کنترل بودن آماره‌ی تعداد نقص‌ها که از توزیع پواسون تبعیت می‌کند بوسیله‌ی نمودار کنترل بررسی می‌شود (برای توضیح بیشتر، به بخش ۲-۲-۳ از [۱] مراجعه کنید). این در حالی است که اهمیت این نقص‌ها از دید مصرف‌کننده و تولیدکننده در عمل معمولاً یکسان نبوده و لذا در چنین مواردی بکارگیری نمودار کنترل $\tilde{C}$ مبتنی بر کیفیت فازی، به جای نمودار کنترل C ، موجه جلوه می‌کند.

برای کنترل تعداد نواقص، ابتدا لازم است تا یک واحد بازرسی تعریف و انتخاب گردد. واحد بازرسی عبارت است از میزان ثابتی از خروجی یک سیستم تولیدی که باید به طور منظم مورد بازرسی قرار گیرد. به طور مثال برای شمارش تعداد نقص‌ها در سیستم‌های تولیدی پیوسته مانند خط تولید رول‌های کاغذ، سیم و یا پارچه، ممکن است واحد بازرسی برابر $4m$ ، $6cm$ و یا $12m^2$ در نظر گرفته شود.

پس از معرفی و تعیین واحد بازرسی، برای رسم نمودار کنترل $\tilde{C}$ بر اساس اطلاعات m واحد بازرسی، کافیت نقاط $(i, \tilde{w}_i)$ را به ازای $i = 1, 2, \dots, m$ بر روی صفحه رسم و به یکدیگر متصل کنیم که در آن

$$\tilde{w}_i = \sum_{j=1}^{w_i} \tilde{N}(x_{ij}) = \sum_{j=1}^{w_i} (1 - \tilde{Q}(x_{ij}))$$

مجموع درجات بی‌کیفیتی (نقص) در واحد بازرسی i ام، مقدار مشخصه کیفیت مربوط به i امین نقص در واحد بازرسی i ام، و w_i تعداد نقص‌ها در واحد بازرسی i ام می‌باشد که در بخش ۲-۲-۳ از [۱] با نماد c_i نشان داده شده است. به این ترتیب، اگر تمامی نقاط رسم شده بین حدود کنترل قرار داشته و هیچ گونه رفتار سیستماتیکی از خود نشان ندهند، آن‌گاه می‌توان تحت کنترل بودن مجموع درجات بی‌کیفیتی‌ها (تعداد نقص‌ها) را نتیجه گرفت و در غیر اینصورت فرایند از کنترل خارج بوده و نیازمند اقدامات اصلاحی می‌باشد.

توجه داشته باشید که برای این نمودار x در فرمول‌های (۲.۳) تعداد نقص‌ها نیست بلکه مشخصه کیفیت (مثلاً مساحت لکه بر بدنه‌ی کالای تولیدی) است. گر چه تابع چگالی احتمال این مشخصه کیفیت برای کاربران معمولاً مجهول است، اما برای ترسیم نمودار کنترل شوهارت

و تشخیص حدود کنترل تنها کافیت مقادیر میانگین و انحراف معیار $\tilde{Q}(X)$ را بر اساس داده‌ها به کمک روابط (۲.۴) برآورد کرد که در آن‌ها $\bar{Q}_i = \frac{\sum_{j=1}^{w_i} \tilde{Q}(x_{ij})}{w_i}$ و $s_i = \sqrt{\frac{\sum_{j=1}^{w_i} [\tilde{Q}(x_{ij}) - \bar{Q}_i]^2}{w_i - 1}}$ به ترتیب برآوردهای میانگین و انحراف معیار در واحد بازرسی i ام هستند. بنابراین با توجه به قضیه ۱.۳، برآورد حدود کنترل و خط مرکزی نمودار $\tilde{C}$ برابر است با

$$\begin{aligned}\widehat{LCL} &= n \left(1 - \bar{Q} - k \frac{s_{\tilde{Q}}}{\sqrt{n}} \right) \\ \widehat{CL} &= n \left(1 - \bar{Q} \right) \\ \widehat{UCL} &= n \left(1 - \bar{Q} + k \frac{s_{\tilde{Q}}}{\sqrt{n}} \right)\end{aligned}\quad (5.4)$$

۵. نتیجه‌گیری

یکی از مفروضات اصلی برای رسم نمودارهای کنترل کیفیت غیرفازی p ، np و C هم‌ارزش و یکسان فرض کردن تمامی معایب/نواقص است و بر اساس همین فرض تنها تعداد معایب/نواقص به عنوان آماره‌ای برای برپایی نمودار کنترل در نظر گرفته می‌شود. ولی در عمل ممکن است با حالت‌هایی مواجه شویم که در آن‌ها برقراری این فرض اولیه غیرمنطقی و نادرست باشد. بدیهی است که در چنین مواقعی، نمودارهای کنترل غیرفازی فوق‌کاری خود را از دست داده و نیاز به روش‌های جایگزین دارند. به عنوان یک روش جایگزین، می‌توان به نمودارهای کنترل کیفیت فازی مطرح شده در این مقاله اشاره کرد که در آن‌ها درجات بی‌کیفیتی‌ها (مقادیر معایب/نواقص) برای تمامی اقلام معیوب/ناقص با یکدیگر مساوی و یکسان در نظر گرفته نمی‌شوند.

مراجع

۱. ع. پرچمی و م. ماشین‌چی، کنترل کیفیت آماری، انتشارات دانشگاه شهید باهنر کرمان، ۱۳۹۱.
۲. ع. پرچمی و م. ماشین‌چی، کیفیت فازی و نسل جدیدی از شاخصهای کارایی، اندیشه آماری، سال دوازدهم، شماره یکم (شماره پیاپی ۱۳)، ۱۳۸۶، ۶۸-۷۶.
۳. س.م. مقدس و م. صالح‌اولیاء، کنترل کیفیت: سیستم، سازماندهی، روش‌های آماری، جهاد دانشگاهی صنعتی شریف، ۱۳۷۰.
۴. د. مونتگمری، کنترل کیفیت آماری، ترجمه‌ی ر. نورالسنا، انتشارات دانشگاه علم و صنعت، ۱۳۷۶.
5. Amirzadeh, V., Mashinchi, M., Parchami, A. (2009), *Construction of p-charts using degree of nonconformity*, Information Sciences 179, 150-160.
6. Amirzadeh, V., Mashinchi, M., Yaghoobi, M.A. (2008), *Construction of control charts using fuzzy multinomial quality*, Journal of Mathematics and Statistics 4, 26-31.
7. Yongting, C. (1996), *Fuzzy quality and analysis on fuzzy probability*, Fuzzy Sets and Systems 83, 283-290.



رگرسیون وزنی فازی کمترین قدرمطلق خطا

جلال چاچی^{۱*} و فائزه ترکیان^۲

^۱دانشگاه سمنان، دانشکده ریاضی، آمار و علوم کامپیوتر، گروه آمار

jchachi@semnan.ac.ir

^۲دانشگاه پیام نور تهران واحد شرق

fa-torkian@yahoo.com

چکیده. در رویکرد کمترین مربعات خطا، برآورد پارامترهای مدل‌های رگرسیون فازی نسبت به داده‌های پرت بسیار حساس می‌باشد. از این جهت معرفی مدل‌های رگرسیون فازی با رویکرد کمترین قدرمطلق خطا می‌تواند جایگزین مناسبی برای مدل‌های رگرسیون فازی با رویکرد کمترین مربعات خطا باشند. در این مقاله رویکرد جدید رگرسیون وزنی فازی کمترین قدرمطلق خطا برای مدل‌سازی مشاهدات ورودی-دقیق و خروجی-فازی معرفی می‌شود که در آن از یک معیار برازش وزنی برای برآورد پارامترهای مدل استفاده می‌شود. بر اساس این معیار وزن بهینه هر مشاهده متناسب با اهمیت و تاثیر آن مشاهده در برآورد پارامترها تعیین می‌شود. مدل پیشنهادی تعمیمی از مدل رگرسیون فازی کمترین قدرمطلق خطا است، زیرا اگر وزن تمام مشاهدات یکسان در نظر گرفته شوند، این مدل تبدیل به یک مدل رگرسیون فازی کمترین قدرمطلق خطای فازی متداول خواهد شد.

۱. مقدمه

در تجزیه و تحلیل مدل‌های رگرسیون فازی (شامل رویکردهای امکانی، رویکردهای مبتنی بر شیوه کمترین مربعات خطا و دیگر رویکردهای تلفیقی)، به منظور برآورد بهتر و دقیق‌تر پارامترها، مجموعه مشاهدات نمونه باید با هدف شناسایی نقاط پرت مورد بررسی و تحلیل قرار گیرد. از طرفی در صورت مشاهده نقطه (یا نقاط) پرت در مجموعه داده‌ها بهتر و مناسب‌تر است که از روش‌های برآوردیابی استوار، به عنوان جایگزین‌های بهتر و مناسب‌تر برای برآورد پارامترها، استفاده شود [۱]. تا کنون روش‌ها و رویکردهای مختلف و متنوعی در برآوردیابی استوار پارامترهای مدل‌های رگرسیون فازی ارائه شده است. این تنوع در روش و رویکرد به دلیل گستردگی در تنوع تعریف مشاهدات پرت در مطالعات مدل‌سازی رگرسیون در محیط فازی است [۳]. در این مقاله به معرفی مدل رگرسیون وزنی فازی با رویکرد کمترین قدرمطلق خطا برای داده‌های ورودی-دقیق و خروجی-فازی پرداخته می‌شود. همچنین فرض می‌شود مجموعه مشاهدات شامل داده‌های نامتعارف و پرت هستند و روش‌های مدل‌سازی مبتنی بر رویکرد کمترین قدرمطلق خطای متداول در مدل‌سازی این داده‌ها برازش مناسبی ارائه نمی‌دهند. روش پیشنهادی بر مبنای یک معیار برازش وزنی مبتنی بر رویکرد کمترین قدرمطلق خطا است و در

واژگان کلیدی. رگرسیون وزنی فازی کمترین قدرمطلق خطا، داده‌های فازی، الگوریتم وزن-دهی تکرار شونده.

*سخنران

حضور مشاهدات پرت، برآوردهای استواری برای پارامترهای مدل ارائه دهد. در این روش مقادیر بهینه پارامترها به همراه وزن‌های بهینه هر مشاهده بر اساس یک الگوریتم وزن-دهی تکرار شونده تعیین می‌شوند. این الگوریتم به‌گونه‌ای عمل می‌کند که داده‌های پرت وزن‌های کوچکتر و داده‌های خوب وزن‌های بزرگتری در برازش مدل دارند. مطالب این مقاله به صورت زیر تدوین شده است. در بخش ۲، برخی از مفاهیم و تعاریف مورد نیاز از مجموعه‌های فازی بیان می‌شوند. در بخش ۳، روش برآوردیابی مدل رگرسیون وزنی فازی بیان می‌شود. در بخش ۴، با حل یک مثال عددی به بررسی مدل رگرسیون فازی کمترین قدرمطلق خطا پرداخته می‌شود. در انتها و در بخش ۵ به بحث و نتیجه‌گیری پرداخته می‌شود.

۲. مفاهیم اولیه

مجموعه فازی $\tilde{A}$ از $\mathbb{R}$ با تابع عضویت $\tilde{A}(x) : \mathbb{R} \rightarrow [0, 1]$ مشخص می‌شود. α -برش مجموعه فازی $\tilde{A}$ برای هر $\alpha \in (0, 1]$ به صورت مجموعه معمولی $A_\alpha = \{x \in \mathbb{R} : \tilde{A}(x) \geq \alpha\}$ تعریف می‌شود و A_0 بستار مجموعه $\{x \in \mathbb{R} : \tilde{A}(x) > 0\}$ است [۵].

تعریف ۱.۲. مجموعه فازی $\tilde{A}$ از $\mathbb{R}$ را یک عدد فازی گوئیم هرگاه برای هر $\alpha \in (0, 1]$ مجموعه‌های $A_\alpha = \{x \in \mathbb{R} : \tilde{A}(x) \geq \alpha\}$ ناتهی، بسته و کراندار باشند.

تعریف ۲.۲. یک رده خاص از اعداد فازی در $\mathbb{R}$ اعداد فازی LR هستند [۵]. یک عدد فازی LR به صورت $\tilde{N} = (n, l, r)_{LR}$ نشان داده می‌شود که در آن $n \in \mathbb{R}$ مرکز، $l \in \mathbb{R}^+$ و $r \in \mathbb{R}^+$ به ترتیب پهناهای چپ و راست عدد فازی و $L : \mathbb{R}^+ \rightarrow [0, 1]$ و $R : \mathbb{R}^+ \rightarrow [0, 1]$ به ترتیب توابع شکل نزولی چپ و راست، با شرط $L(0) = R(0) = 1$ هستند. تابع عضویت α -برش عدد فازی $\tilde{N} = (n, l, r)_{LR}$ به صورت زیر تعریف می‌شود

$$\tilde{N}(x) = \begin{cases} L\left(\frac{n-x}{l}\right) & \text{if } x \leq n, \\ R\left(\frac{x-n}{r}\right) & \text{if } x \geq n. \end{cases}$$

$$N_\alpha = [n - L^{-1}(\alpha)l, n + R^{-1}(\alpha)r], \quad \alpha \in [0, 1].$$

نوع خاصی از اعداد فازی LR ، اعداد فازی مثلثی هستند. تابع عضویت α -برش عدد فازی مثلثی $\tilde{N} = (n, l, r)_T$ به صورت زیر تعریف می‌شود

$$\tilde{N}(x) = \begin{cases} \frac{x-(n-l)}{l} & \text{if } x \in [n-l, n], \\ \frac{(n+r)-x}{r} & \text{if } x \in (n, n+r]. \end{cases}$$

$$N_\alpha = [n - (1-\alpha)l, n + (1-\alpha)r], \quad \alpha \in [0, 1].$$

تعریف ۳.۲. فرض کنید $\tilde{N} = (n, l_n, r_n)_{LR}$ و $\tilde{M} = (m, l_m, r_m)_{LR}$ دو عدد فازی LR باشند و $\lambda \in \mathbb{R} - \{0\}$ یک عدد حقیقی باشد. در این صورت حساب اعداد فازی به صورت زیر

تعریف می‌شود

$$\begin{aligned}\lambda \otimes \widetilde{M} &= \begin{cases} (\lambda m, \lambda l_m, \lambda r_m)_{LR} & \text{if } \lambda > 0, \\ (\lambda m, |\lambda| r_m, |\lambda| l_m)_{RL} & \text{if } \lambda < 0, \end{cases} \\ \widetilde{M} \oplus \widetilde{N} &= (m+n, l_m+l_n, r_m+r_n)_{LR}.\end{aligned}$$

۳. رگرسیون وزنی فازی کمترین قدرمطلق خطا

در این بخش، روش جدید رگرسیون وزنی فازی کمترین قدرمطلق خطا برای مدل‌سازی داده‌های

$$(\widetilde{y}_1, x_1), \dots, (\widetilde{y}_n, x_n), \quad (1.3)$$

معرفی می‌شود. در این داده‌ها $\widetilde{y}_i = (y_i, l_i, r_i)_{LR}$ و x_i به ترتیب i -امین مشاهدات متغیرهای خروجی-فازی و ورودی-دقیق هستند. بدون اینکه از کلیت مساله کاسته شود فرض می‌شود که مقادیر متغیر ورودی نامنفی هستند. برای برازش یک مدل رگرسیون فازی به مجموعه مشاهدات (۱.۳)، ابتدا روابط زیر بین متغیرهای ورودی-خروجی در نظر گرفته می‌شوند

$$\begin{aligned}(y, g(l), g(r))_{LR} &= \widetilde{\beta}_0 \oplus (\widetilde{\beta}_1 \otimes x) \\ &= (\beta_0, \sigma_0, \theta_0)_{LR} \oplus ((\beta_1, \sigma_1, \theta_1)_{LR} \otimes x) \\ &= (\beta_0 + \beta_1 x, \sigma_0 + \sigma_1 x, \theta_0 + \theta_1 x)_{LR},\end{aligned} \quad (2.3)$$

که در آن $g: \mathbb{R}^+ \rightarrow \mathbb{R}^+$ تابعی معکوس‌پذیر است. با در نظر گرفتن تابع g در روابط بالا مقادیر پهنای برآورد شده برای متغیر خروجی فازی همواره نامنفی خواهد بود. در نهایت پس از برآورد بهینه پارامترهای β_0 و β_1 و با توجه به معکوس‌پذیری تابع g ، مدل رگرسیون فازی در برازش به مجموعه مشاهدات (۱.۳) به صورت زیر حاصل می‌شود

$$\begin{aligned}\widetilde{y} &= (y, l, r)_{LR} \\ &= (\beta_0 + \beta_1 x, g^{-1}\{\sigma_0 + \sigma_1 x\}, g^{-1}\{\theta_0 + \theta_1 x\})_{LR}.\end{aligned} \quad (3.3)$$

در مدل‌سازی مجموعه مشاهدات (۱.۳) مطابق مدل رگرسیون فازی (۳.۳)، پارامترهای $\widetilde{\beta}_0$ و $\widetilde{\beta}_1$ از کمینه کردن تابع برازش وزنی زیر به دست می‌آیند

$$E = \sum_{i=1}^n w_i e_i, \quad (4.3)$$

که در آن w_i وزن مشاهده i ام است و e_i خطای مشاهده i ام، به صورت زیر است

$$\begin{aligned}e_i &= \mathcal{D}\left((y_i, g(l_i), g(r_i))_{LR}, (\widetilde{\beta}_0 \oplus (\widetilde{\beta}_1 \otimes x_i))\right) \\ &= |y_i - (\beta_0 + \beta_1 x_i)| + c_1 |g(l_i) - (\sigma_0 + \sigma_1 x_i)| \\ &\quad + c_2 |g(r_i) - (\theta_0 + \theta_1 x_i)|.\end{aligned}$$

برای برآورد بهینه پارامترهای $\widetilde{\beta}_0$ و $\widetilde{\beta}_1$ بر اساس تابع برازش وزنی (۴.۳) از الگوریتم وزن-دهی تکرار شونده زیر استفاده می‌کنیم. زیرا در تابع برازش وزنی (۴.۳) وزن‌ها وابسته به خطاها،

خطاها وابسته به مقادیر پارامترهایی هستند که خود از تابع هدفی به دست می‌آیند که وابسته به وزن‌ها است.

الگوریتم: محاسبه برآوردهای وزنی برای پارامترهای مدل رگرسیون فازی (۳.۳).

ورودی الگوریتم: مشاهدات متغیرهای ورودی-خروجی.

خروجی الگوریتم: مقدار بهینه برآورد شده پارامترهای مدل رگرسیون فازی (۳.۳)،

$\hat{\beta} = [\hat{\beta}_0, \hat{\beta}_1]$ به همراه وزن‌های بهینه $w = [w_1, \dots, w_n]$
گام ۵: با استفاده از وزن‌های اولیه زیر برای مشاهده i ام

$$w_i^0 = \frac{1}{\max(h_{ii}, \frac{2}{n})}, \quad i = 1, \dots, n,$$

که در آن h_{ii} درایه قطری ماتریس $H = X(X'X)^{-1}X'$ است، تابع هدف زیر را کمینه نمایید

$$\begin{aligned} E &= \sum_{i=1}^n w_i^0 e_i \\ &= \sum_{i=1}^n w_i^0 \left[|y_i - (\beta_0 + \beta_1 x_i)| + c_1 |g(l_i) - (\sigma_0 + \sigma_1 x_i)| \right. \\ &\quad \left. + c_2 |g(r_i) - (\theta_0 + \theta_1 x_i)| \right]. \end{aligned}$$

جواب مساله بهینه‌سازی بالا را $\hat{\beta}^0 = [\hat{\beta}_0^0, \hat{\beta}_1^0]$ رار دهید.

گام t : برای تکرارهای $t = 1, 2, \dots$ تا همگرایی مقادیر پارامترهای بهینه، مراحل زیر را انجام دهید:

(۱) خطای مدل را بر اساس مقدار برآورد شده $\hat{\beta}^{t-1}$ در تکرار $t-1$ به صورت زیر محاسبه نمایید

$$\begin{aligned} e^{t-1} &= [e_1^{t-1}, \dots, e_n^{t-1}], \\ e_i^{t-1} &= \mathcal{D} \left((y_i, g(l_i), g(r_i)), \hat{\beta}_0^{t-1} \oplus (\hat{\beta}_1^{t-1} \otimes x_i) \right), \end{aligned}$$

(۲) وزن‌های جدید زیر را محاسبه نمایید

$$\begin{aligned} w_i^t &= \frac{|1 - h_{ii}|}{\max\{e_i^{t-1}, M_{e^{t-1}}\}}, \quad i = 1, \dots, n, \\ M_{e^{t-1}} &= \text{Median}\{e_1^{t-1}, \dots, e_n^{t-1}\}. \end{aligned}$$

جدول ۱: مجموعه داده‌های ورودی-حقیقی و خروجی-فازی به همراه وزن‌های بهینه برآورد شده هر مشاهده

$\hat{w}_i$	x_i	$\tilde{y}_i$	i
0/2423	2/00	$(5/8, 3/5, 4/2)_T$	1
0/0915	0/00	$(0/8, 0/5, 0/2)_T$	2
0/0722	1/13	$(13/9, 8/5, 7/0)_T$	3
0/2423	2/00	$(4/0, 2/4, 3/0)_T$	4
0/2432	2/19	$(1/6, 1/0, 1/0)_T$	5
0/1383	0/25	$(1/5, 0/9, 1/2)_T$	6
0/1827	0/75	$(8/2, 5/0, 5/0)_T$	7
0/2089	4/25	$(1/8, 1/1, 2/0)_T$	8
0/0133	8/50	$(0/5, 0/5, 1/0)_T$	9

(۳) تابع هدف زیر را کمینه نمایید

$$\begin{aligned}
 E &= \sum_{i=1}^n w_i^t e_i \\
 &= \sum_{i=1}^n w_i^t \left[|y_i - (\beta_0 + \beta_1 x_i)| + c_1 |g(l_i) - (\sigma_0 + \sigma_1 x_i)| \right. \\
 &\quad \left. + c_2 |g(r_i) - (\theta_0 + \theta_1 x_i)| \right].
 \end{aligned}$$

جواب مساله بهینه‌سازی بالا را $\hat{\beta}^t = [\hat{\beta}_0^t, \hat{\beta}_1^t]$ قرار دهید.

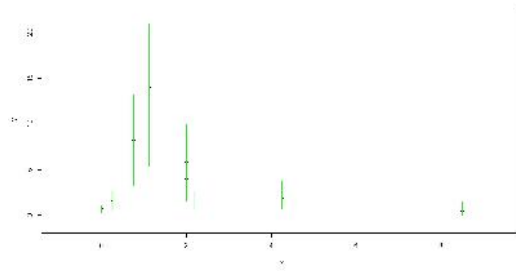
در نهایت پس از خاتمه الگوریتم بالا (همگرایی مقادیر پارامترهای بهینه) مدل رگرسیون وزنی فازی به صورت زیر نوشته می‌شود

$$\begin{aligned}
 \hat{\tilde{y}} &= (\hat{y}, \hat{l}, \hat{r})_{LR} \\
 &= (\hat{\beta}_0 + \hat{\beta}_1 x, g^{-1}\{\hat{\sigma}_0 + \hat{\sigma}_1 x\}, g^{-1}\{\hat{\theta}_0 + \hat{\theta}_1 x\})_{LR}.
 \end{aligned}$$

نکته ۱.۳. در الگوریتم بالا، نقاط اهرمی و نقاط با مقادیر خطای بزرگ (e_i بزرگ) وزن‌های کوچکتری می‌گیرند. همچنین اگر در این الگوریتم وزن کلیه مشاهدات یکسان در نظر گرفته شود، مدل رگرسیون وزنی کمترین قدرمطلق خطا تبدیل به مدل رگرسیون کمترین قدرمطلق خطای متداول می‌شود.

۴. مثال عددی

در این بخش، توانایی و عملکرد مدل پیشنهاد شده در این مقاله در شناسایی و کم اثر کردن مشاهده (یا مشاهدات) پرت در برآوردیابی پارامترهای مدل توضیح داده می‌شود. برای این منظور داده‌های جدول ۱ را در نظر بگیرید. نمودار پراکنش مشاهدات متغیر ورودی در مقابل ابهام (یا



شکل ۱: نمودار پراکنش مشاهدات متغیر ورودی در مقابل ابهام (یا پهنا) متغیر خروجی.

پهنا) متغیر خروجی در شکل ۱ نشان داده شده است. طبق اطلاعات این نمودار، واضح است که مشاهده نهم، یعنی $(x_9, \tilde{y}_9) = (8/50, (0/5, 0/5, 1/0)_T)$ یک مشاهده پرت در این مجموعه داده است. با اعمال الگوریتم بیان شده در این مقاله، مدل رگرسیون وزنی فازی کمترین قدرمطلق خطا برای این مجموعه داده به صورت زیر به دست می‌آید

$$\hat{y}_{WLA} = (5.07 - 0.53x, \exp\{1.35 - 0.24x\}, \exp\{1.43 - 0.10x\})_{LR}.$$

همچنین وزن‌های بهینه این مشاهدات $\hat{w}_i$ ، $i = 1, \dots, 9$ ، که میزان تاثیر هر مشاهده در برآورد پارامترهای مدل را نشان می‌دهد، در ستون چهارم جدول ۱ گزارش شده است. همانطور که انتظار می‌رفت، مشاهده نهم کوچکترین وزن را در بین دیگر مشاهدات دارد، زیرا مقدار x بزرگی دارد. همچنین بر اساس وزن بهینه برآورد شده برای هر مشاهده، بدیهی است که داده‌های دوم و سوم رتبه‌های بعدی را در نامتعارف بودن دارند.

۵. نتیجه‌گیری

مدل‌های رگرسیون فازی کمترین مربعات حساسیت زیادی نسبت به داده‌های پرت دارند. از این جهت، در این مطالعه برآوردهای رگرسیونی وزنی فازی با رویکرد کمترین قدرمطلق خطا در مدل‌سازی داده‌های فازی معرفی شد که نسبت به مشاهدات پرت حساس نمی‌باشد. این روش تعمیمی از رویکرد رگرسیون فازی کمترین قدرمطلق خطای متداول است و همچنین توانایی شناسایی و کم اثر ساختن مشاهدات پرت را در برآورد بهینه پارامترهای مدل دارد. به عبارتی در رویکرد معرفی شده، وزن هر مشاهده متناسب با تاثیر مطلوبی که در برآوردیابی پارامترها دارد مشخص می‌شود. بنابراین، به مشاهدات پرت، که اغلب تاثیر نامطلوبی در برآورد پارامترها دارند، وزن کمتری در برآوردیابی پارامترها اختصاص داده می‌شود.

در مطالعات آتی، کاربرد عملی روش پیشنهاد شده در این مقاله در مدل‌سازی داده‌های واقعی قابل توجه است. همچنین تاثیر انواع مختلف مشاهدات پرت در برآورد بهینه پارامترها و برازش این مدل‌ها به مجموعه داده‌ها قابل تحلیل و بررسی است.

مراجع

1. Andersen, R. (2007), *Modern Methods for Robust Regression*, Sage, Thousand Oaks, CA.
2. Chachi, J., Taheri, S.M., Fattahi, S. and Hosseini Ravandi, S.A. (2017), *Two robust fuzzy regression models and their applications to predict imperfections of cotton yarn*, Journal of Textiles and Polymers, In Press.
3. D'Urso, P., Massari, R. and Santoro, A. (2011), *Robust fuzzy regression analysis*, Inform. Sci. 181, 4154–4174.
4. Kelkinnama, M. and Taheri, S.M. (2012), *Fuzzy least-absolute regression using shape preserving operations*, Inform. Sci. 214, 105–120.
5. Zimmermann, H.J. (2001), *Fuzzy Set Theory and Its Applications*, 4th edn, Kluwer Nihoff: Boston.



کاربرد رگرسیون فازی در مدل‌بندی شاخص‌های کیفی آب

جلال چاچی^{۱*}، مریم میرزایی^۲، و بهنام عرب حسن آبادی^۳

^۱دانشگاه سمنان، دانشکده ریاضی، آمار و علوم کامپیوتر، گروه آمار
Jchachi@semnan.ac.ir

^۲دانشگاه سمنان، دانشکده ریاضی، آمار و علوم کامپیوتر، گروه آمار
Maryam.mirzayi69@gmail.com

^۳دانشگاه خوارزمی، گروه ریاضی
Behnam.arab1992@gmail.com

چکیده. کیفیت آب یک حوضه بستگی به عوامل متعددی مانند زمین‌شناسی، سیلاب‌ها، مدیریت کشت و غیره دارد. جامدات، ذرات معلق یا مواد محلول موجود در آب (یا فاضلاب) بر کیفیت آب آشامیدنی، آب کشاورزی و آب مورد استفاده در صنعت تأثیر نامطلوبی می‌گذارند. هدف اصلی این مطالعه بررسی تعیین مدل بهینه برآورد کیفیت آب بر اساس یک مدل رگرسیون فازی مبتنی بر انتخاب مجموعه‌ای از موثرترین متغیرهای مستقل تأثیرگذار در کیفیت آب است.

۱. مقدمه

در یک مدل رگرسیونی انتخاب متغیرها از اهمیت زیادی برخوردار است. این که چه متغیرهایی باید وارد مدل شوند و چه متغیرهایی باید از مدل حذف شوند، نکته مهمی است که همواره باید به آن توجه شود. در برازش یک مدل رگرسیونی مناسب به دلیل هزینه زیادی که همراه با تهیه اطلاعات بر مبنای تعداد زیاد متغیرها وجود دارد، باید تا حد امکان تعداد کمی از متغیرها در نظر گرفته شود. از طرفی دیگر برای افزایش دقت و اطمینان از پیش‌بینی مدل رگرسیونی انتخاب بهینه تعداد متغیرها امری ضروری و اساسی است. بنابراین برای برازش بهترین مدل رگرسیونی باید مناسب‌ترین و مفیدترین متغیرها انتخاب شوند. لذا در رگرسیون چندگانه فازی، یک موضوع مهم، تعریف معیار انتخاب متغیر برای انتخاب مناسب مجموعه‌هایی از متغیرهای توضیحی می‌باشد. برای انتخاب بهترین مدل معیارهایی وجود دارد که هرکدام از این معیارهای انتخاب متغیر در رگرسیون فازی نیز می‌توانند مورد بررسی و تحلیل قرار گیرند. در چارچوب رگرسیون معمولی روش‌های مختلفی در این زمینه پیشنهاد شده است [۱]. در ادامه به معرفی روش ضریب تعیین برای انتخاب متغیر در مدل رگرسیون فازی تعریف و بررسی می‌شود. از طرفی کاربرد این رویکرد را در مدل‌سازی داده‌های واقعی در بررسی کیفیت آب (آب‌شناسی) مورد بررسی و مطالعه قرار می‌دهیم.

واژگان کلیدی. رگرسیون فازی، انتخاب متغیر در مدل رگرسیون فازی، داده‌های کیفی آب.
* سخنران

۲. مفاهیم اولیه

مجموعه فازی $\tilde{A}$ از $\mathbb{R}$ با تابع عضویت $\tilde{A}(x) : \mathbb{R} \rightarrow [0, 1]$ مشخص می‌شود. α -برش مجموعه فازی $\tilde{A}$ برای هر $\alpha \in (0, 1]$ به صورت مجموعه معمولی $A_\alpha = \{x \in \mathbb{R} : \tilde{A}(x) \geq \alpha\}$ تعریف می‌شود و A_0 بستار مجموعه $\{x \in \mathbb{R} : \tilde{A}(x) > 0\}$ است [۴].

تعریف ۱.۲. مجموعه فازی $\tilde{A}$ از $\mathbb{R}$ را یک عدد فازی گوییم هرگاه برای هر $\alpha \in (0, 1]$ مجموعه‌های $A_\alpha = \{x \in \mathbb{R} : \tilde{A}(x) \geq \alpha\}$ ناتهی، بسته و کراندار باشند.

تعریف ۲.۲. یک رده خاص از اعداد فازی در $\mathbb{R}$ اعداد فازی LR هستند [۴]. یک عدد فازی LR به صورت $\tilde{N} = (n, l, r)_{LR}$ نشان داده می‌شود که در آن $n \in \mathbb{R}$ مرکز، $l \in \mathbb{R}^+$ و $r \in \mathbb{R}^+$ به ترتیب پهناهای چپ و راست عدد فازی و $L : \mathbb{R}^+ \rightarrow [0, 1]$ و $R : \mathbb{R}^+ \rightarrow [0, 1]$ به ترتیب توابع شکل نزولی چپ و راست، با شرط $L(0) = R(0) = 1$ هستند. تابع عضویت α -برش عدد فازی $\tilde{N} = (n, l, r)_{LR}$ به صورت زیر تعریف می‌شود

$$\tilde{N}(x) = \begin{cases} L\left(\frac{n-x}{l}\right) & \text{if } x \leq n, \\ R\left(\frac{x-n}{r}\right) & \text{if } x \geq n. \end{cases}$$

$$N_\alpha = [n - L^{-1}(\alpha)l, n + R^{-1}(\alpha)r], \quad \alpha \in [0, 1].$$

نوع خاصی از اعداد فازی LR ، اعداد فازی مثلثی هستند. تابع عضویت α -برش عدد فازی مثلثی $\tilde{N} = (n, l, r)_T$ به صورت زیر تعریف می‌شود

$$\tilde{N}(x) = \begin{cases} \frac{x-(n-l)}{l} & \text{if } x \in [n-l, n], \\ \frac{(n+r)-x}{r} & \text{if } x \in (n, n+r]. \end{cases}$$

$$N_\alpha = [n - (1-\alpha)l, n + (1-\alpha)r], \quad \alpha \in [0, 1].$$

تعریف ۳.۲. فرض کنید $\tilde{N} = (n, l_n, r_n)_{LR}$ و $\tilde{M} = (m, l_m, r_m)_{LR}$ دو عدد فازی LR باشند و $\lambda \in \mathbb{R} - \{0\}$ یک عدد حقیقی باشد. در این صورت حساب اعداد فازی به صورت زیر تعریف می‌شود

$$\lambda \otimes \tilde{M} = \begin{cases} (\lambda m, \lambda l_m, \lambda r_m)_{LR} & \text{if } \lambda > 0, \\ (\lambda m, |\lambda| r_m, |\lambda| l_m)_{RL} & \text{if } \lambda < 0, \end{cases}$$

$$\tilde{M} \oplus \tilde{N} = (m+n, l_m+l_n, r_m+r_n)_{LR}.$$

۳. رگرسیون فازی کمترین مربعات

در این بخش با استفاده از رویکرد کمترین مربعات به مدل‌سازی مجموعه مشاهدات زیر می‌پردازیم

$$(\tilde{y}_1, \mathbf{x}_1), \dots, (\tilde{y}_n, \mathbf{x}_n).$$

این مشاهدات مربوط به یک یا چند متغیر مستقل و یک متغیر وابسته می‌باشد. در این مشاهدات:

(۱) $x_i = (x_{0i}, x_{1i}, \dots, x_{ki})$ ، i -امین مشاهدات متغیرهای مستقل است،
 (۲) $\tilde{y}_i = (y_i, l, r)_{LR}$ ، $(i = 1, \dots, n, k < n, x_{0i} = 1)$ ، i -امین مشاهدات متغیر وابسته است.

اکنون به منظور مدل سازی داده های بیان شده، مدل رگرسیون فازی زیر در نظر گرفته می شود

$$\begin{aligned}(y, G(l), G(r))_{LR} &= \tilde{\beta}_0 \oplus (\tilde{\beta}_1 \otimes x_1) \oplus \dots \oplus (\tilde{\beta}_k \otimes x_k) \\ &= (\beta_0, \gamma_0, \delta_0)_{LR} \oplus ((\beta_1, \gamma_1, \delta_1)_{LR} \otimes x_1) \oplus \dots \\ &\quad \dots \oplus ((\beta_k, \gamma_k, \delta_k)_{LR} \otimes x_k) \\ &= (\beta_0 + \beta_1 x_1 + \dots + \beta_k x_k, \gamma_0 + \gamma_1 x_1 + \dots + \gamma_k x_k, \\ &\quad \delta_0 + \delta_1 x_1 + \dots + \delta_k x_k)_{LR}.\end{aligned}$$

که در آن $G : (0, \infty) \rightarrow \mathbb{R}$ تابعی معکوس پذیر است [۲]. چون تابع $G : (0, \infty) \rightarrow \mathbb{R}$ معکوس پذیر است، لذا متغیر وابسته به صورت زیر برآورد می شود

$$\begin{aligned}(\hat{y}, \hat{l}, \hat{r}) &= (\hat{\beta}_0 + \hat{\beta}_1 x_1 + \dots + \hat{\beta}_k x_k, G^{-1}(\hat{\gamma}_0 + \hat{\gamma}_1 x_1 + \dots + \hat{\gamma}_k x_k), \\ &\quad G^{-1}(\hat{\delta}_0 + \hat{\delta}_1 x_1 + \dots + \hat{\delta}_k x_k))_{LR}.\end{aligned}$$

در ادامه پارامترهای مدل بالا یعنی

$$\tilde{\beta} = [\tilde{\beta}_0, \tilde{\beta}_1, \dots, \tilde{\beta}_k]$$

که در آن

$$\tilde{\beta}_j = (\beta_j, \gamma_j, \delta_j)_{LR}, \quad j = 0, 1, \dots, k$$

با استفاده از رویکرد کمترین مربعات از مساله بهینه سازی زیر به دست می آیند

$$\begin{aligned}\min \quad & SSE = SSE_1 + SSE_2 + SSE_3 \\ \text{s.t.} \quad & \beta_j \in \mathbb{R}, \quad \gamma_j \in \mathbb{R}, \quad \delta_j \in \mathbb{R}, \quad j = 0, 1, \dots, k\end{aligned}$$

که در آن

$$\begin{aligned}SSE_1 &= \sum_{i=1}^n [y_i - (\beta_0 + \beta_1 x_{1i} + \dots + \beta_k x_{ki})]^2 \\ SSE_2 &= \sum_{i=1}^n [G(l_i) - (\gamma_0 + \gamma_1 x_{1i} + \dots + \gamma_k x_{ki})]^2 \\ SSE_3 &= \sum_{i=1}^n [G(r_i) - (\delta_0 + \delta_1 x_{1i} + \dots + \delta_k x_{ki})]^2\end{aligned}$$

در نهایت پس از برآورد بهینه پارامترها، مدل بهینه رگرسیون فازی به صورت زیر نوشته می شود

$$\begin{aligned}\hat{y} &= (\hat{y}, \hat{l}, \hat{r})_{LR} \\ &= (\hat{\beta}_0 + \hat{\beta}_1 x_1 + \dots + \hat{\beta}_k x_k, G^{-1}(\hat{\gamma}_0 + \hat{\gamma}_1 x_1 + \dots + \hat{\gamma}_k x_k), \\ &\quad G^{-1}(\hat{\delta}_0 + \hat{\delta}_1 x_1 + \dots + \hat{\delta}_k x_k))_{LR}.\end{aligned}$$

۴. معیار انتخاب متغیر در مدل رگرسیون فازی

در رگرسیون چندگانه فازی، یک موضوع مهم، تعریف معیار انتخاب متغیر برای انتخاب مناسب مجموعه‌هایی از متغیرهای توضیحی می‌باشد. در چارچوب رگرسیون معمولی روش‌های مختلفی در این زمینه پیشنهاد شده است. در ادامه به معرفی روش ضریب تعیین برای انتخاب متغیر در مدل رگرسیون فازی پرداخته می‌شود. برای مطالعه بیشتر در این زمینه به مراجع [۱، ۲، ۳] مراجعه نمایید. ضریب تعیین نشان می‌دهد که مدل چه میزان از تغییرات متغیر خروجی فازی را توسط متغیرهای مستقل توضیح می‌دهد. اکنون مدل زیر را در نظر بگیرید

$$(y, G(l), G(r))_{LR} = \tilde{\beta}_0 \oplus (\tilde{\beta}_1 \otimes x_1) \oplus \dots \oplus (\tilde{\beta}_k \otimes x_k)$$

برای این مدل، ضریب تعیین تعدیل شده به صورت زیر محاسبه می‌شود

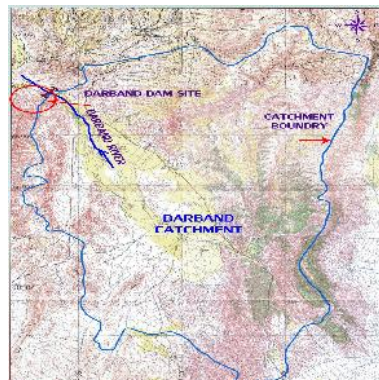
$$R_{adj}^2 = 1 - \frac{n-1}{n-p} \frac{SSE}{SST}$$

که در آن

$$\begin{aligned} SST &= \sum_{i=1}^n D^2(\tilde{y}_i, \bar{\tilde{y}}_i) \\ &= \sum_{i=1}^n (y_i - \bar{y})^2 + \frac{1}{3} \sum_{i=1}^n (G(l_i) - \bar{G(l)})^2 + \frac{1}{3} \sum_{i=1}^n (G(r_i) - \bar{G(r)})^2 \\ &= SST_1 + SST_2 + SST_3 \\ SSE &= \sum_{i=1}^n D^2(\tilde{y}_i, \hat{\tilde{y}}_i) \\ &= \sum_{i=1}^n (y_i - \hat{y}_i)^2 + \frac{1}{3} \sum_{i=1}^n (G(l_i) - \widehat{G(l_i)})^2 + \frac{1}{3} \sum_{i=1}^n (G(r_i) - \widehat{G(r_i)})^2 \\ &= SSE_1 + SSE_2 + SSE_3 \\ \tilde{y}_i &= \frac{1}{n} \oplus_{i=1}^n (y_i, G(l_i), G(r_i))_{LR} \\ &= (\frac{1}{n} \sum_{i=1}^n y_i, \frac{1}{n} \sum_{i=1}^n G(l_i), \frac{1}{n} \sum_{i=1}^n G(r_i))_{LR} \\ &= (\bar{y}, \bar{G(l)}, \bar{G(r)})_{LR}. \end{aligned}$$

عملکرد انتخاب متغیر بر مبنای ضریب تعیین R_{adj}^2 ابتدا مجموعه‌های $A_1, A_2, \dots, A_m$ را به صورت زیر در نظر می‌گیریم ($m = 2^{k+1} - 2$)

- (۱) مجموعه A_1 شامل معادله‌های رگرسیون فازی که دارای یک متغیر مستقل هستند.
- (۲) مجموعه A_2 شامل معادله‌های رگرسیون فازی که دارای دو متغیر مستقل هستند.
- (۳) مجموعه A_3 شامل معادله‌های رگرسیون فازی که دارای سه متغیر مستقل هستند.
- (۴) مجموعه A_k شامل معادله‌های رگرسیون فازی که دارای k متغیر مستقل هستند.



شکل ۱: مختصات محل و شکل جغرافیایی حوزه آبریز دربند سملقان در خراسان شمالی به طول جغرافیایی $56^{\circ} 58' 48''$ و عرض جغرافیایی $37^{\circ} 36' 21''$.



شکل ۲: شکل جغرافیایی و مسیر دسترسی به حوزه آبریز دربند سملقان.

آنگاه درون هر مجموعه، معادله‌ای که دارای بیشترین مقدار R_{adj}^2 است، انتخاب می‌شود (در صورت وجود دو یا چند معادله با R_{adj}^2 های یکسان که از بقیه بزرگتر هستند، هر دو یا چند معادله انتخاب می‌شوند). تمامی معادله‌های انتخابی بر حسب مقادیر R_{adj}^2 مرتب می‌شوند. در انتها آن مجموعه‌ای که بیشترین مقدار R_{adj}^2 را دارد و اختلاف معنی‌داری با مجموعه‌های دیگر دارد، به عنوان بهترین معادله رگرسیون انتخاب می‌شود [۱]. توجه کنید که در این تصمیم‌گیری ممکن است معیارهایی مانند هزینه، زمان، افزایش برازش معنی‌دار برای مدل و عوامل محیطی دیگر دخالت داشته باشند.

۵. مطالعه کاربردی

۱۰۵. معرفی مکان مطالعه. حوزه آبریز دربند سملقان با مساحت ۱۱۱۷ کیلومتر مربع و محیط ۱۸۱/۸ کیلومتر در خراسان شمالی قرار دارد (شکل ۱). ارتفاع کوتاهترین، متوسط و

بلندترین نقطه حوضه به ترتیب: 68° ، $1285/8$ و 2725 متر از سطح دریاست. مختصات خروجی حوضه به طول جغرافیایی $37^{\circ} 36' 21''$ و عرض جغرافیایی $56^{\circ} 58' 48''$ است (شکل ۲). این مقاله قسمتی از پروژه مطالعات سد مخزنی بر روی رودخانه دربند (خراسان شمالی) است که توسط شرکت مهندسی مشاور تحقیقات خاک مهاراب مطالعه شده است (شکل های ۱ و ۲). در این مطالعه داده های کیفی آب طی آمارهای طولانی مدت از دبی رودخانه و نمونه های کیفی آب رودخانه به تعداد زیادی برای متغیرهای مختلفی ثبت شده است. نتایج و اطلاعات کیفی برخی از این نمونه ها برای متغیرهایی که در بخش آتی تعریف می شوند، مورد بررسی و تحلیل قرار گرفته اند.

۲.۵. متغیر وابسته: میزان جامدات محلول در آب. متغیر Tds^1 مهمترین شاخص برای بررسی کیفیت آب است که ذرات جامد محلول در آب یا به عبارتی کل نمک های محلول یا مقدار باقی مانده خشک را اندازه گیری می کند و واحد اندازه گیری آن میلی گرم در لیتر mg/lit و یا قسمت در میلیون ppm می باشد. کل نمک های محلول یا مقدار باقی مانده خشک را می توان در صورتی که آب آبیاری فاقد بی کربنات باشد به سادگی با تبخیر حجم معینی از آب و اندازه گیری وزن املاح باقی مانده تعیین کرد. در آب هایی که مقدار قابل توجهی بی کربنات وجود دارد با استفاده از این روش ممکن است تنها نصف وزن بی کربنات اندازه گیری شود، زیرا حدود نیمی از بی کربنات در اثر حرارت از بین رفته و مقدار Tds بدست آمده با مقدار واقعی آن متفاوت خواهد بود. لذا در آزمایشگاه های اندازه گیری میزان واقعی اندازه گیری شده این شاخص ممکن است مقداری کمتر یا بیشتر از میزان واقعی کل نمک های محلول یا مقدار باقی مانده خشک را ثبت نماید. لذا به نظر می رسد با استفاده از نظریه مجموعه های فازی بتوان رویکردی در جمع آوری و ثبت این مشاهدات در یک محیط نادقیق ارائه داد. برای این منظور مشاهدات این متغیر به صورت اعداد فازی مثلثی مقدار (شامل یک مقدار مرکزی و یک مقدار پهنای چپ و راست) ثبت شده اند (جدول ۱).

۳.۵. متغیرهای مستقل. در کیفیت آب آشامیدنی متغیرهای مختلف و متنوعی تاثیرگذار هستند، که به برخی از آنها در ادامه اشاره می شود. مشاهدات مربوط به این متغیرها در جدول ۱ گزارش شده است.

(۱) x_1 : متغیر $Debi$ بیانگر میزان مقدار آب جاری در یک رودخانه در یک مکان خاص است و واحد اندازه گیری آن متر مکعب بر ثانیه می باشد.

(۲) x_2 : متغیر Ec (هدایت الکتریکی) نشان دهنده میزان املاح و شوری موجود در آب می باشد. در واقع تعیین غلظت املاح محلول در آب را نشان می دهد که واحد اندازه گیری آن میکروموس بر سانتی متر می باشد. هدایت الکتریکی آب نشان دهنده میزان املاح هادی موجود در آب می باشد و در واقع یکی از راه های ساده تعیین غلظت املاح محلول در آب، اندازه گیری هدایت الکتریکی است. آب مقطر یا آب خالص تقریباً هادی جریان الکتریسیته نیست ولی اگر در آب نمک های محلول وجود داشته باشد، آب را هادی جریان الکتریسیته می کند.

(۳) x_3 : متغیر PH میزان غلظت اسیدی یا بازی بودن آب را نشان می دهد.

¹Total dissolved solid

جدول ۱: مشاهدات متغیرهای مورد مطالعه

$tds = (c, l, r)_T$	$sakhti$	cl	na	ph	ec	$debi$
(980, 210, 210) _T	3	5.2	7.0	7.8	1556	0.61
(822, 150, 150) _T	2.65	4.2	6.4	7.9	1305	0.53
(1017, 250, 250) _T	2.55	6.2	8.2	8	1614	0.44
$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$
(764, 128, 128) _T	2.93	3.8	6.9	7.7	1250	0.62
(316, 87, 87) _T	1.48	2.05	0.8	8.4	520	2.72

(۴) x_4 : متغیر Na میزان سدیم موجود در آب را نشان می‌دهد.

(۵) x_5 : متغیر Cl میزان کلر موجود در آب را نشان می‌دهد. کلر با مولکول‌های اکسیژن و هیدروژن آب واکنش داده و اسید هیپوکلریک را تشکیل می‌دهد که واحد اندازه‌گیری آن ppm است. کلر مایع از نمک‌ها تولید شده، به همین خاطر شور است. کلر برای ضد عفونی کننده آب مورد استفاده قرار می‌گیرد. کلر گاز سمی و بسیار واکنش‌پذیر است. از کاربرد آن در تصفیه آب برای از بین بردن باکتری‌ها و سایر میکروب‌های موجود در ذخایر آب آشامیدنی اشاره کرد. همچنین کلر در آب‌های داخل خاک و رسوبات موجود است. کلر مایع از نمک‌ها تولید می‌شود، به همین دلیل ذاتا شور است. در صورت افزایش آن به آب موجب افزایش درجه سختی آب نیز می‌شود.

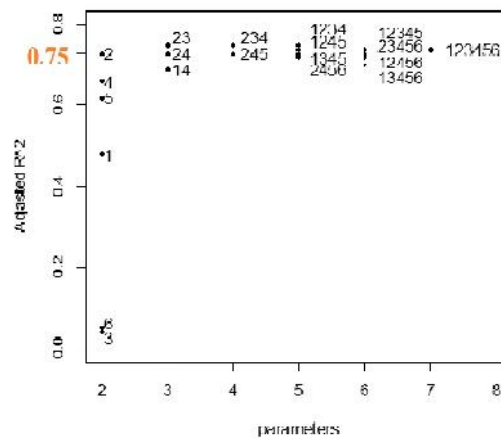
(۶) x_6 : متغیر $Sakhti$ سختی آب را اندازه‌گیری می‌کند که بر حسب واکنش‌های یون‌های کلسیم، منیزیم، کربنات و بی‌کربنات کلسیم با یکدیگر می‌باشد.

۴.۵. نتایج عددی. با اعمال رویکرد مطرح شده در این مقاله، نتایج بصورت اطلاعات نمایش داده شده در شکل ۳ گزارش می‌شود. در این نمودار مقادیر ضریب تعیین برای مدل‌های مختلف (مدل‌های با تعداد متغیرهای مستقل متفاوت) نشان داده شده است. با توجه به رویکرد معرفی شده، چون اختلاف معنی‌داری بین مدل با متغیر x_2 و دیگر مدل‌ها با تعداد متغیرهای متفاوت وجود ندارد، لذا این مدل به عنوان مدل بهینه انتخاب می‌شود که بصورت زیر است

$$\hat{y} = (0/73 + 0/65x_2, \exp\{4/75 + 0/0006x_2\}, \exp\{4/75 + 0/0006x_2\}).$$

۶. نتیجه‌گیری

مطالعات بسیاری با رویکرد تلفیق روش‌های کلاسیک آماری و روش‌های فازی، در زمینه رگرسیون فازی انجام گرفته است. رگرسیون فازی از لحاظ مقایسه با رگرسیون آماری (از نظر قابلیت‌ها، کاربردها، مزیت‌ها) نیز مورد توجه است و تنوع بسیاری دارد. تاکنون رویکردهای مختلفی در زمینه رگرسیون فازی مورد مطالعه و بررسی بسیاری از محققین قرار گرفته است. اما در زمینه‌های جانبی مربوط به این رویکردها، مانند استنباط درباره پارامترهای مدل رگرسیون فازی، همبستگی بین متغیرهای بکار گرفته شده در یک مدل رگرسیون فازی و انتخاب بهینه متغیرهای مدل رگرسیون فازی مطالعات گسترده‌ای صورت نگرفته است. لذا در این مطالعه



شکل ۳: نمودار مقادیر مختلف ضریب تعیین برای مدل‌های مختلف با تعداد متغیرهای مختلف.

به بررسی و تحلیل رویکرد انتخاب بهینه متغیرهای مدل رگرسیون فازی پرداخته شد. مشابه رویکردهای انتخاب متغیر در مدل‌های رگرسیون کلاسیک، در اینجا نیز انجام این کار یک روش منحصر به فرد نمی‌باشد. به عبارتی، در کاربرد این موضوع که چه متغیرهایی باید انتخاب شوند، نیاز به تعریف معیارهایی دور از قضاوت‌های شخصی دارد. بنابراین، انتخاب بهترین مدل، نیازمند تعیین یک معیار انتخاب مناسب برای تصمیم‌گیری است. معیار انتخاب، شاخصی است که برای هر کدام از مدل‌ها با انتخابی از متغیرهای مستقل محاسبه می‌شود و سپس به مقایسه مدل‌ها می‌پردازد. به این ترتیب می‌توان مدل‌های مورد نظر را به ترتیب از بهترین تا بدترین مدل در نظر گرفت و بهترین مدل را انتخاب کرد. برای انتخاب بهترین مدل معیارهای مختلف متنوعی را می‌توان در تحقیقات آتی تعریف نمود. اما این نکته را باید همواره در نظر داشت که ”هنگامی که این روش‌ها برای یک مسئله به کار گرفته می‌شوند، لزوماً به پاسخ‌های یکسان منجر نخواهند شد که این موضوع باعث اختلال خواهد شد، هرچند برای بسیاری از مسائل، پاسخ یکسانی خواهد داشت.“

مراجع

۱. ع. بازرگان لاری، رگرسیون خطی کاربردی، انتشارات دانشگاه شیراز، ۱۳۸۱.
2. D'Urso P. and Massari, R. (2013), *Weighted least squares and least median squares estimation for the fuzzy linear regression analysis*, Metron 71, 279–306.
3. D'Urso P. and Santoro A. (2006), *Goodness of fit and variable selection in the fuzzy multiple linear regression*, Fuzzy Sets and Systems 157, 2627–2647.
4. Zimmermann, H.J. (2001), *Fuzzy Set Theory and Its Applications*, 4th edn, Kluwer Nihoff: Boston.



منطق فازی، احتمال، استنتاج منطقی و احتمالات

سید محمد امین خاتمی^۱ *

^۱ دانشگاه صنعتی بیرجند، گروه علوم کامپیوتر
khatami@birjandut.ac.ir, khatami@ipm.ir

چکیده. در مقاله مروری که پیش رو دارید، گذری کوتاه خواهیم داشت بر مفهومی مجموعه‌های فازی و احتمال و با ذکر چند مثال تفاوت مفهومی که بین این دو وجود دارد را بهتر روشن می‌کنیم. سپس به بیان چند نکته در باب استفاده از استنتاج در منطق کلاسیک و منطق فازی و احتمال می‌پردازیم. در انتها به صورت مختصر منطق احتمالاتی را معرفی می‌کنیم.

۱. پیش‌گفتار

مفهوم مجموعه‌های فازی در سال ۱۹۶۵ توسط آقای لطفی علی‌عسکرزاده، پرفسور ایرانی‌الاصل دانشکده برق دانشگاه کالیفرنیا معرفی شد. ایشان چند سال بعد از معرفی مجموعه‌های فازی عبارت منطق فازی را مورد استفاده قرار دادند و سپس در مقالات بعدی خود این مفهوم را نیز به صورت دقیق‌تر تعریف کردند. به گفته خود پرفسور زاده [۳]، گرچه انتقاداتی به این دو مفهوم بنیادی از همان ابتدا وارد شد، و افراد زیادی مقابل این مفاهیم بدلیل مسئله تشکیکی که در آنها نهفته است موضع گرفتند، اما بدلیل نیاز گسترده جامعه علمی روز به این مفاهیم، این مفاهیم مورد استقبال طیف گسترده‌ای از جامعه علمی و بخصوص متخصصانی که با کاربردهای روزمره علم سروکار دارند قرار گرفت. از طرف دیگر، نظریه احتمالات که در بدو ابداع مفهوم مجموعه‌های فازی عمر زیادی داشت نیز بنوعی با مفاهیم مبهم سروکار داشت. این باعث شد افراد زیادی مفهوم مجموعه‌ها و منطق فازی را با نظریه احتمالات اشتباه کنند و حتی یکسان در نظر بگیرند تا جایی که منطق فازی را بنوعی همان نظریه احتمالات در لباسی دیگر بدانند. ما در قسمت اول این مقاله با ذکر چند مثال، تفاوت اساسی که بین این دو مفهوم وجود دارد را توضیح خواهیم داد. موضوع دیگری که موقع استفاده از منطق فازی مهم است، توجه به مبانی منطقی و خواص انواع مختلف این منطق‌ها است. متأسفانه موقعی که یک محقق به تنهایی از منطق فازی برای انجام یک کار استفاده می‌کند، عدم آشنایی وی با مبانی منطق مورد استفاده، ممکن است موجبات تضعیف یک ایده خیلی قوی را فراهم آورد، در حالی که اندکی توجه به خواص منطق مورد استفاده و شیوه‌های استنتاج در منطق مورد استفاده وی، ممکن است مسیر انجام دادن کار را برای او

واژگان کلیدی. مجموعه‌های فازی، منطق فازی، احتمال.
* سخنران

روشن‌تر کند. از طرفی گرچه سیستم‌های اصل موضوعی مختلفی برای منطق‌های فازی مختلف وجود دارد اما حداقل منطق‌های فازی مبتنی بر نرم مثلثی پیوسته خیلی نزدیک هم هستند [۲] و لذا مطالعه این خواص، حداقل در مورد منطق‌های فازی مبتنی بر نرم مثلثی پیوسته زمان زیادی لازم ندارد.

۲. مجموعه‌های فازی و نظریه احتمالات

معمولاً منطق ریاضی را به عنوان ابزاری برای استنتاج منطقی از روی چند واقعیت (گزاره) مشخص می‌شناسیم. در منطق کلاسیک این استدلال با توجه به معنی‌شناسی مبتنی بر درست و غلط بودن واقعیت‌ها انجام می‌شود و در منطق‌های چند مقداری و منطق فازی این استدلال مبتنی بر ارزش گزاره‌ها که محدود به درست و غلط محض نیستند و بین این دو هم می‌توانند باشند، انجام می‌شود. اما چیزی که در مورد همه منطق‌های مذکور، چه منطق کلاسیک، چه منطق‌های چند مقداری و چه منطق فازی در استدلال کردن وجود دارد این واقعیت است که

«هر گزاره ارزش معلومی از درستی دارد که از روی ارزش گزاره‌های اتمیک تشکیل دهنده آن و توابع درستی روابط منطقی تشکیل دهنده آن گزاره بدست می‌آید».

توجه کنید که در اینجا واژه «معلوم» با نفس کلمه «ابهام» که در فازی بودن نهفته است، تناقضی ندارد. در واقع وقتی می‌خواهیم برای گزاره «علی قد بلند است»، درجه‌ای از درستی تعیین کنیم، با توجه به درجات مختلف قد و داشتن اندازه قد علی، باید این درجه درستی را معلوم کنیم. درجه‌ای از درستی که بین «درست محض» و «غلط محض» قرار دارد و «معلوم» است که چه درجه‌ای است.

در نظریه احتمالات هم ارزش هر گزاره‌ای با یک احتمال خاص و با توجه به پیشامدهای ممکن، مشخص می‌گردد، اما تفاوت اساسی این ارزش احتمالی با ارزش فازی در این است که «ارزش احتمالی» قبل از اندازه گرفتن قد علی و با ماهیتی «اعتقادی» و با تکیه بر شواهد مشخص می‌گردد، اما برای تعیین «ارزش فازی» ممکن است قد اندازه گرفته شده باشد.

از طرف دیگر در بیشتر منطق‌های فازی، وقتی ارزش دو گزاره φ و ψ معلوم باشد، به ازای هر رابط منطقی Δ ، ارزش گزاره $\varphi \Delta \psi$ نیز با توجه به تابع درستی رابط منطقی Δ و ارزش دو گزاره φ و ψ قابل محاسبه است. این در حالی است که در نظریه احتمالات با دانستن ارزش احتمالی دو گزاره φ و ψ ، ارزش احتمالی گزاره $\varphi \Delta \psi$ در حالت کلی قابل تعیین نیست.

تفاوت سوم نظریه احتمالات و منطق فازی در نحوه استنتاج کردن است. در منطق کلاسیک و بیشتر منطق‌های فازی، استنتاج مستقل از زمان می‌باشد. به عبارت دیگر وقتی اطلاعات ما برای استنتاج کردن افزایش یافت، این اطلاعات جدید به مجموعه مفروضات ولی بدون در نظر گرفتن هیچ ترتیبی اضافه می‌شوند. در حالت طبیعی انتظار می‌رود مفروضات باید دارای یک ترتیب باشند و این فرض جدید به عنوان آخرین فرض به مجموعه مفروضات اضافه شود ولی در استنتاج با منطق کلاسیک و بیشتر منطق‌های فازی اینطور نیست و هیچ ترتیبی روی مجموعه مفروضات قائل نمی‌شوند. این در حالی است که در استنتاج‌های آماری زمان اهمیت دارد. مثلاً

در قاعده احتمال شرطی، احتمال وقوع یک گزاره بشرط این که گزاره دیگری رخ داده باشد را محاسبه می‌کنند که نشان دهنده اهمیت فاکتور زمان می‌باشد. چهارمین تفاوت هم موقع تشکیل مجموعه‌های فازی بر پایه شواهد دیده می‌شود. شاید این تفاوت به دو تعبیر از اصل موضوع گسترش^۱ در نظریه مجموعه‌ها مربوط باشد. این اصل که در واقع اولین اصل از اصول موضوعه زرمولو-فرانکلین^۲ برای نظریه مجموعه‌ها است، در شکل اصلی‌اش (نحوی) بیان می‌کند که دو مجموعه وقتی مساوی هستند که خواص درونی یکسانی داشته باشند:

دو مجموعه وقتی مساوی هستند که عناصر یکسانی داشته باشند. به عبارت دیگر

$$\forall A \forall B (A = B \leftrightarrow \forall c (c \in A \leftrightarrow c \in B))$$

شکل معنی‌شناختی اصل توسیع که به اصل موضوع مصداقیت در فلسفه مشهور است، بیان می‌کند دو مجموعه وقتی مساوی هستند که بروندهای یکسانی را نتیجه دهند:

دو مجموعه وقتی مساوی هستند که زیرمجموعه‌های مجموعه‌های یکسانی باشند. به عبارت دیگر

$$\forall A \forall B (A = B \leftrightarrow \forall C (A \subseteq C \leftrightarrow B \subseteq C))$$

به نظر می‌رسد در بیان نحوی تساوی، بنوعی مجموعه‌های فازی و نظریه احتمالات دو مفهوم یکسان هستند اما در بیان معنی‌شناختی تساوی، مجموعه‌های فازی بیشتر خودنمایی می‌کنند و احتمالات کمتر دیده می‌شود.

فرض کنید بخواهند در مورد اینکه شخص مریضی مبتلا به یک بیماری خاص است تصمیم بگیرند. بدیهی است که در این صورت به بررسی آثار آن بیماری در شخص می‌پردازند. بعد از بررسی آثار، چند بیماری A_1 و A_2 و A_3 و ... محتمل به نظر می‌رسند. از دیدگاه نحوی در مورد این شخص، بعد از مشاهده نتایج و بررسی سوابق آثار بیماری در بیماران دیگر، درجه‌ای از احتمال به بیماری خواهیم داشت که یک زیرمجموعه فازی به صورت $\{\frac{A_1}{.3}, \frac{A_2}{.4}, \frac{A_3}{.1}, \frac{A_4}{.2}, \frac{A_5}{.1}, \dots\}$ است. این مجموعه فازی که منشاء آن از نظریه احتمالات است نشان می‌دهد 0.3 افرادی که نشانه‌های این بیماری را دارند مبتلا به بیماری A_1 بوده‌اند و قس علی‌هذا. اما در دیدگاه معنی‌شناختی بیشتر از اینکه مشاهدات سوابق بیماران قبلی مهم باشد، مشاهده آثار بیماری مهم است که حتی بعضی از این آثار هنوز آنقدر برای بشر شناخته شده نیست که به آنها توجه کند. در اینجا بررسی مرادفات فرد با اجتماع، سوابق شغلی، سوابق خانوادگی و ... نیز اهمیت پیدا می‌کنند. در این حالت ارزشی که به بیماری A_n منسوب می‌شود در واقع نشان دهنده این است که در این بیماری چه درصدی از آثار شناخته شده ظهور پیدا می‌کند و یک زیر مجموعه فازی به صورت $\{\frac{A_1}{.25}, \frac{A_2}{.18}, \frac{A_3}{.25}, \frac{A_4}{.12}, \frac{A_5}{.1}, \dots\}$ بدست می‌آید که شاید تعبیری احتمالی از آن نتوان ارائه کرد.

^۱ Axiom of extensionality
^۲ ZFC

۳. استنتاج منطقی و احتمالاتی

در این بخش اشاره‌ای خواهیم داشت به خواص استنتاج در منطق کلاسیک و منطق فازی و سپس معرفی خواهیم داشت از منطق احتمالاتی. یکی از بهترین صورت‌های منطق کلاسیک و منطق‌های فازی که هم برای استنتاج در ریاضیات و سایر علوم مناسب است و هم بررسی خواص مورد علاقه منطق‌دانان، برای آن پیچیدگی خیلی زیادی ندارد منطق مرتبه اول است. در این صورت از منطق، گزاره‌ها به کمک یک «زبان مرتبه اول» که شامل تعدادی ثابت، تابع و محمول می‌باشد، ساخته می‌شوند. سپس با توجه به یک «جهان» یا «عالم سخن» که عناصر زبانی در آن تعبیر می‌شوند، ارزش گزاره‌ها قابل تعیین کردن است. تا مرحله نوشتن گزاره‌ها از روی زبان تنها فرق انواع مختلف منطق‌های مرتبه اول، در روابط منطقی و سورهای مورد استفاده می‌باشد. تفاوت اصلی بین منطق کلاسیک با منطق‌های فازی در معنی‌شناسی می‌باشد.

معنی‌شناسی یک منطق در واقع با تعبیر عناصر زبانی در عالم سخن شروع می‌شود و سپس با توجه به توابع درستی مربوط به روابط منطقی و سورها، معنای تمام گزاره‌ها قابل تعیین کردن است. در معنی‌شناسی عناصر زبانی، تعبیر ثوابت و توابع در همه منطق‌ها یکسان است و این تعبیر محمول‌ها، روابط منطقی و سورها است که باعث ایجاد تفاوت بین منطق کلاسیک با منطق‌های فازی می‌شود.

علم منطق بعد از معنی‌شناسی، به مطالعه مفهوم استنتاج و اینکه چطور یک نتیجه از روی مجموعه‌ای از مفروضات در یک منطق خاص بدست می‌آید، می‌پردازد. در اینجا حساب‌های منطقی مختلفی برای مفهوم «اثبات» بوجود می‌آیند. اما در عین حال در همه این حساب‌های منطقی، هدف اصلی این است که اولاً

درستی^۳ : هر چیزی که اثبات می‌شود به لحاظ معنی‌شناسی درست باشد،

و ثانیاً

تمامیت^۴ : هر چیزی که به لحاظ معنی‌شناسی درست است قابل اثبات باشد.

معمولاً حساب منطقی همیشه طوری چیده می‌شود که درستی برقرار باشد. در واقع اصولی که در استنتاج از آنها استفاده می‌شود باید درست باشند. اما تمامیت شاید همواره برقرار نباشد. در منطق کلاسیک مرتبه اول، تمامیت نیز برقرار است. علاوه بر این در منطق‌های فازی مبتنی بر نرم مثلثی پیوسته نیز نگارش‌هایی از تمامیت قابل اثبات است [۲]. اما برای یک منطق فازی دلخواه، شاید هنوز حساب منطقی خوبی که تمامیت نیز در آن درست باشد، طراحی نشده باشد. به نظر می‌رسد توجه به این موضوع و سایر جنبه‌های منطقی که در اینجا مجال برای بیان آنها نیست، افق‌های دید دیگری را در اختیار محققین قرار دهد.

تا اینجا اشاره‌ای داشتیم به اهمیت نگاه منطقی در سایر علوم. مثالی که در زیر مطرح می‌کنیم بنوعی نشان می‌دهد روش استنتاجی که در مسائل کاربردی به آن نیاز داریم فراتر از چیزی است که منطق ریاضی و منطق فازی مطرح می‌کنند، و حتی محدود کردن استنتاج به روش‌های موجود

^۳ Soundness

^۴ Completeness

در منطق، ممکن است مانع منطقی بودن بشود! مثال زیر بنوعی نشان می‌دهد که روش‌های استنباطی احتمالات هم باید به شیوه‌های استنتاج منطق ریاضی افزوده شود.

پارادکس مونته‌هال^۵ : فرض کنید که در یک مسابقه تلویزیونی شرکت کرده‌اید و میان سه در باید یکی را انتخاب کنید. پشت یکی از درها یک ماشین است و پشت دو در دیگر دو بز. شما یکی از درها را انتخاب می‌کنید (مثلاً در شماره یک). مجری برنامه که می‌داند پشت هر در چه چیزی است، در دیگری را باز می‌کند (مثلاً در شماره سه) و به شما نشان می‌دهد که پشتش یک بز است. بعد از شما می‌پرسد که می‌خواهید در شماره یک را با شماره دو تاخت بزنید؟ آیا به سود شماست که انتخابتان را عوض کنید؟

در ظاهر به نظر می‌رسد که سودی در عوض کردن انتخاب وجود ندارد. در واقع از هر ۱۰ نفر ۹ نفر معتقدند که عوض کردن انتخاب سودی ندارد. طراح مسئله آقای «واس سوانت^۶» بعد از یک بحث طولانی با «پل اردوش^۷» یکی از برجسته‌ترین ریاضیدانان تاریخ نمی‌تواند او را قانع کند که عوض کردن انتخاب اولیه شانس بردن جایزه را به میزان زیادی افزایش می‌دهد و اردوش حتی تا وقتی شبیه سازی کامپیوتری را ندیده بود، اعتقاد داشت که عوض کردن انتخاب سودی ندارد! اما واقعیت این است که با عوض کردن انتخاب، شانس بردن ماشین $\frac{2}{3}$ خواهد شد. روش‌های مختلفی مانند شبیه‌سازی کامپیوتری، استدلال استقرایی با توجه به آزمایش واقعی و استفاده از قاعده احتمال شرطی، نشان می‌دهند که عوض کردن انتخاب اولیه شانس بردن جایزه را افزایش می‌دهند. اشکال کار در چیست که به زعم عده‌ای زیادی از مردم و حتی دانشمندان، مسئله به شکل پارادکس نمایان می‌شود؟ شاید عادت کردن افراد به روش‌های استدلال کلاسیکی که از طفولیت با آنها انس گرفته‌اند، باعث آن باشد. یکی دیگر از سوالات بحث برانگیز که در منطق کلاسیک مطرح می‌شود این است که چرا از تناقض می‌توان هر نتیجه‌ای گرفت؟ این موضوع بر خلاف شهود ما است و طبیعی به نظر نمی‌رسد! آیا تابع درستی روابط منطقی در منطق کلاسیک اشکال دارند؟ آیا تناقض را باید باور داشت و از آن هر چیزی را نتیجه نگرفت؟

دو مسئله فوق و مسائل دیگری که در اینجا فرصت مطرح کردن آنها نیست و این واقعیت که در بیشتر استدلال‌ها، مفروضات به طور حتمی ممکن است درست نباشند، باعث می‌شود که به منطقی فراتر از منطق کلاسیک فکر کنیم. توجه کنید در این حال دو رویکرد ذیل بیشتر به چشم می‌آیند:

- یک رویکرد در نظر گرفتن درجه‌ای از درستی برای مفروضات است که همان منطق فازی می‌باشد،
- و رویکرد دوم در نظر گرفتن درجه احتمال درستی برای مفروضات می‌باشد

اما منطق فازی دو مسئله اول را حل نمی‌کند. در مورد هر دو مسئله مطرح شده، بیشتر منطق‌های فازی جوابی شبیه منطق کلاسیک و یا نزدیک به منطق کلاسیک ارائه می‌دهند. بعلاوه توجه داشته

^۵ Monty Hall problem

^۶ vos Savant

^۷ Paul Erdos

باشید که در مورد پارادکس مانتی‌هال، یا ماشین پشت درب خاصی قرار دارد و یا قرار ندارد و اینطور نیست که قسمتی از آن پشت یک درب باشد. لذا منطق فازی ظاهراً مناسب نیست. اینجا به نظر می‌رسد رویکرد دوم حرفی برای گفتن داشته باشد. در واقع ترکیب ظرفیت نظریه احتمالات برای مهار کردن عدم قطعیتی که در منطق استقرایی وجود دارد، موجب پیدایش صورت‌گرایی‌های غنی‌تر و گویاتری از منطق کلاسیک و منطق‌های فازی شده است. در ادامه یکی از طرح‌های پیشنهادی اخیر [۱] برای گسترش احتمالاتی منطق کلاسیک گزاره‌ای را به طور مختصر معرفی می‌کنیم. ابتدا چند تعریف مقدماتی از منطق کلاسیک را بیان می‌کنیم.

تعریف ۱.۳. فرض کنید φ و ψ دو گزاره و Σ یک تئوری (مجموعه‌ای از گزاره‌ها) باشد.

- φ را توتولوژی می‌نامند هرگاه تحت هر نوع ارزش‌دهی گزاره‌های اتمیک، ارزش درست داشته باشد.
- می‌گویند φ منطقاً معادل ψ است هرگاه $\varphi \leftrightarrow \psi$ توتولوژی باشد.
- می‌گویند φ منطقاً ψ را ایجاب می‌کند هرگاه $\varphi \rightarrow \psi$ توتولوژی باشد.
- می‌گویند Σ منطقاً ψ را ایجاب می‌کند هرگاه تحت هر ارزش‌دهی که همه گزاره‌های Σ درست باشند، ارزش ψ نیز درست باشد.
- می‌گویند φ ناسازگار است هرگاه منطقاً تناقض $(A \wedge \neg A)$ را ایجاب کند.

اکنون تعریف کلیدی منطق احتمالاتی^۸ که به هر گزاره‌ای، احتمالی از وقوع را نسبت می‌دهد به شکل زیر است.

تعریف ۲.۳. اگر تابع p که روی مجموعه گزاره‌ها تعریف شده است خواص زیر را داشته باشد، آن را یک تابع احتمال می‌نامند:

- (۱K) به ازای هر گزاره φ ، $0 \leq p(\varphi) \leq 1$.
- (۲K) به ازای هر گزاره φ که توتولوژی باشد، $p(\varphi) = 1$.
- (۳K) به ازای هر دو گزاره φ و ψ ، اگر φ منطقاً ψ را ایجاب کند آنگاه $p(\varphi) \leq p(\psi)$.
- (۴K) به ازای هر دو گزاره ناسازگار φ و ψ ، $p(\varphi \vee \psi) = p(\varphi) + p(\psi)$.

خواص فوق به اصول موضوعه کولموگوروف^۹ مشهورند. در زیر لیستی از بعضی از قضایای کاربردی که می‌توان با توجه به خواص استنتاجی منطق گزاره‌ای و اصول موضوعه کولموگوروف بدست آورد را می‌بینید:

- (۱P) به ازای هر دو گزاره φ و ψ ، $p(\varphi) \leq p(\varphi \vee \psi)$.
- (۲P) به ازای هر گزاره φ ، $p(\neg\varphi) = 1 - p(\varphi)$.
- (۳P) به ازای هر گزاره φ ، اگر φ منطقاً غلط باشد، آنگاه $p(\varphi) = 0$.
- (۴P) به ازای هر دو گزاره منطقاً معادل φ و ψ ، $p(\varphi) = p(\psi)$.
- (۵P) به ازای هر دو گزاره φ و ψ ، $p(\varphi) + p(\psi) = p(\varphi \wedge \psi) + p(\varphi \vee \psi)$.
- (۶P) به ازای هر دنباله دو به دو ناسازگار از گزاره‌ها مثل $\{\varphi_i\}_{i=1}^n$ ،

$$p\left(\bigvee_{i=1}^n \varphi_i\right) = \sum_{i=1}^n p(\varphi_i)$$

(NP) به ازای هر دو گزاره φ و ψ ، $p(\varphi \wedge \psi) \geq p(\varphi) + p(\psi) - 1$

(NP) اگر از $\{\varphi_i\}_{i=1}^n$ بتوان φ را نتیجه گرفت و علاوه $1 = \min_{1 \leq i \leq n} p(\varphi_i)$ ، آنگاه $p(\varphi) = 1$. همچنین اگر نتوان از $\{\varphi_i\}_{i=1}^n$ گزاره φ را نتیجه گرفت، تابع احتمالی مثل q روی مجموعه گزاره‌ها وجود خواهد داشت که $1 = \min_{1 \leq i \leq n} q(\varphi_i)$ ولی $q(\varphi) = 0$.

اگر تابع عدم قطعیت u از روی تابع احتمال p به صورت $u(\varphi) = 1 - p(\varphi)$ تعریف شود آنگاه بعضی از قضایای کاربردی‌تر مربوط به عدم قطعیت عبارتند از:

(U) به ازای هر دو گزاره φ و ψ ، اگر φ منطقاً ψ را نتیجه دهد، آنگاه $u(\psi) \leq u(\varphi)$.
 علاوه اگر ψ از روی φ نتیجه نشود، آنگاه تابع عدم قطعیتی مثل v وجود خواهد داشت که $v(\varphi) = 0$ ولی $v(\psi) = 1$.

(U) به ازای هر گزاره منطقاً درست φ ، $u(\varphi) = 0$ و به ازای هر گزاره منطقاً غلط φ ، $u(\varphi) = 1$.

(U) به ازای هر دنباله $\{\varphi_i\}_{i=1}^n$ از گزاره‌ها، $u\left(\bigwedge_{i=1}^n \varphi_i\right) \leq \sum_{i=1}^n u(\varphi_i)$ و لذا اگر دنباله گزاره‌های فوق ناسازگار باشد، آنگاه $1 \geq \sum_{i=1}^n u(\varphi_i)$.

مشکل یکی از مسائلی که قبلاً در مورد آن بحث کردیم، کمی در اینجا حل می‌شود. با توجه به (U) و (U) براحتی می‌توان دید اگرچه هر گزاره‌ای مثل φ را می‌توان از تناقض نتیجه گرفت، ولی هیچ امیدی برای اینکه عدم قطعیت چنین گزاره‌ای نزدیک صفر باشد نداریم و فقط می‌دانیم عدم قطعیت آن کمتر یا مساوی یک خواهد بود.

در مورد پارادکس مونتی‌هال و سایر مسائلی در حوزه منطق که به نوعی نیازمند ترتیبی روی گزاره‌ها به لحاظ وقوع زمانی آنها هستیم، آدامز قاعده احتمال شرطی را پیشنهاد می‌دهد که ظاهراً مفید است. اگر p یک تابع احتمال و χ یک گزاره باشد، آنگاه تابع احتمال p_χ روی هر گزاره φ به صورت

$$p_\chi(\varphi) = \begin{cases} \frac{p(\varphi \wedge \chi)}{p(\chi)} & p(\chi) > 0 \\ 1 & p(\chi) = 0 \end{cases}$$

تعریف می‌شود^{۱۰} که آن را با نماد $p(\varphi|\chi)$ و یا گاهی اوقات با نماد $p(\chi \Rightarrow \varphi)$ نمایش می‌دهند. توجه کنید که $\varphi \Rightarrow \chi$ یک گزاره نیست و این فقط یک نمادگذاری برای احتمال وقوع گزاره φ به شرط وقوع گزاره ψ می‌باشد. همانطور که قبلاً گفتیم در منطق گزاره‌ای و لذا در گزاره

^{۱۰} به ازای هر گزاره φ ، $0 \leq p_\chi(\varphi) \leq 1$. اگر $\varphi \vdash \chi$ آنگاه $p_\chi(\varphi) = 1$. اگر $\{\chi, \varphi\} \vdash \psi$ ، آنگاه $p_\chi(\psi) \leq p_\chi(\varphi)$. اگر $p(\chi) > 0$ و $\{\chi \wedge \varphi, \chi \wedge \psi\} \vdash \perp$ ، آنگاه $p_\chi(\varphi \vee \psi) = p_\chi(\varphi) + p_\chi(\psi)$. در اینجا $A \vdash B$ به این معنی است که $A \rightarrow B$ توتولوژی است.

شرطی $\varphi \rightarrow \chi$ که معادل منطقی گزاره $\neg \chi \vee \varphi$ است، وقوع زمانی گزاره‌ها در نظر گرفته نمی‌شود و همه مفروضات در یک زمان اتفاق می‌افتند. به لحاظ مقدار تابع احتمال هم

$$\begin{aligned} p(\chi \rightarrow \varphi) &= 1 - p(\chi \wedge \neg \varphi) \\ &= 1 - p(\chi)p(\neg \varphi|\chi) \\ &= 1 - p(\chi)(1 - p(\varphi|\chi)) \\ &= 1 - p(\chi)(1 - p(\chi \Rightarrow \varphi)) \end{aligned}$$

مثلاً اگر در مورد گزاره

اگر کارت اول تک باشد آنگاه کارت دوم هم تک خواهد بود

در یک مجموعه ۵۲ تایی که ۴ تا تک دارد، بخواهیم ارزش احتمالی را محاسبه کنیم، داریم:

$$\begin{aligned} p(\chi \rightarrow \varphi) &= 1 - (1/13) \times (1 - 3/51) \simeq 0.928 \bullet \\ p(\chi \Rightarrow \varphi) &= 3/51 \simeq 0.059 \bullet \end{aligned}$$

که نشان می‌دهد هنگام وقوع یک گزاره قبل از گزاره دیگر تابع احتمال p و در نتیجه $p(\chi \rightarrow \varphi)$ مناسب نیست و باید از تابع احتمال p_χ یعنی $p(\chi \Rightarrow \varphi)$ که در زمان دیگری است استفاده شود.

۴. نکات نهایی

با توجه به مطالب بخش دوم مقاله، منطق فازی و نظریه احتمالات دو مفهوم کاملاً متمایزند که هر دو بنوعی در مورد عدم قطعیت، اطلاعاتی ارائه می‌دهند. بنابراین در جاهای مختلف باید توجه داشت که کدامیک برای منظور مورد استفاده مناسب‌تر است. در بخش سوم دو موضوع مهم که به نوعی در مقابل هم هستند را مطرح کردیم. اولین موضوع مربوط به توجه به جنبه‌های منطقی منطق فازی مورد استفاده توسط محققین است. در واقع این توجه ممکن است موجب بهبود نتایج تحقیق گردد. موضوع دوم هم برعکس به مهم بودن توجه منطق‌دانان به سایر علوم و به خصوص نظریه احتمالات و اطلاعات تاکید دارد. با ذکر پارادکس مانتی هال، دیدیم گاهی منطق مورد استفاده ما ممکن است مانع منطقی بودن ما بشود. در انتهای بخش سوم هم یکی از رویکردهای ترکیب ظرفیت نظریه احتمالات با منطق کلاسیک را معرفی کردیم، که توضیحی برای بعضی مسائل مناقشه برانگیز منطق کلاسیک ارائه می‌دهد.

مراجع

1. Adams, E. W. (1996). *A primer of probability logic*, University of Chicago Press.
2. Hajek, P. (1998). *Metamathematics of fuzzy logic*, Springer Science.
3. Zadeh, L. A. (2015). *Fuzzy logic, a personal perspective*, Fuzzy Sets and Systems, 281, 4-20.



نگاهی دیگر به معنی‌شناسی در منطق فازی

سید محمد امین خاتمی^۱ * و قاسم خاکشور^۲

^۱ دانشگاه صنعتی بیرجند، گروه علوم کامپیوتر

khatami@birjandut.ac.ir, khatami@ipm.ir

^۲ دانشگاه فنی و حرفه‌ای، آموزشکده فنی و حرفه‌ای دختران بیرجند

ghkhakshor@gmail.com

چکیده. در مقاله حاضر بعد از مرور معنی‌شناسی متداول منطق‌های چندمقداری و منطق‌های فازی، دوگان این معنی‌شناسی را معرفی می‌کنیم. این معنی‌شناسی جدید، نوع نگاه دیگری را با خود در مطالعه منطق‌های چندمقداری و منطق‌های فازی مبتنی بر نرم مثلث پیوسته به همراه دارد. در واقع در انتهای مقاله نشان می‌دهیم این معنی‌شناسی منجر به پیدایش یک متریک بین گزاره‌ها می‌شود.

۱. پیش‌گفتار

اولین منطق چندمقداری در سال ۱۹۲۰ توسط لوکاسویچ^۱ معرفی شد. لوکاسویچ با بیان اینکه «اگر اصل طرد شق ثالث^۲ را در منطق قبول نکنیم چه اتفاقی خواهد افتاد» و سپس در نظر گرفتن یک مقدار درستی دیگر علاوه بر «درست» و «غلط» به مجموعه مقادیر درستی، منطق لوکاسویچ^۳ مقداری را معرفی کرد. با توجه به فرض بنیادی لوکاسویچ مبني بر رد اصل طرد شق ثالث، منطق‌های چندمقداری فقط به لحاظ مجموعه مقادیر درستی با منطق کلاسیک فرق دارند و تفاوت دیگری ندارند. به طور دقیق‌تر در منطق‌های چندمقداری شبیه منطق کلاسیک ارزش همه گزاره‌ها از روی تابع درستی مربوط به روابط منطقی و ارزش گزاره‌های اتمیک قابل محاسبه است. کار لوکاسویچ در سال ۱۹۳۰ توسط خود وی و تارسکی^۳ به یک منطق «بی‌نهایت مقداری شمارا» تعمیم داده شد و تا به امروز نگارش‌های متعدد دیگری از منطق‌های چندمقداری با مجموعه‌های مقادیر درستی متفاوت و معنی‌شناسی مبتنی بر توابع درستی متفاوت معرفی شده‌اند و بنا بر نیازهای مختلف هنوز در حال معرفی هستند. به عنوان مثال، منطق‌های فازی یکی از حالت‌های خاص منطق‌های چند مقداری هستند که در آنها بازه واحد بسته $[0, 1]$ به عنوان مجموعه مقادیر

واژگان کلیدی. منطق‌های چندمقداری، منطق فازی، معنی‌شناسی.

* سخنران

^۱ Łukasiewicz

^۲ principle of excluded middle

^۳ Tarski

درستی در نظر گرفته شده است. در عین حال باید توجه داشت که در بیشتر اوقات، مجموعه مقادیر درستی انتخاب شده یک شبکه با کوچکترین و بزرگترین عنصر می باشد. موضوع مهمی که معمولاً موقع به کار بردن یک منطق چندمقداری مورد غفلت واقع می شود، توجه به مبانی منطقی و خواص منطقی منطق مورد استفاده است. البته این موضوع تا جایی که به مبانی پایه ای منطق آسیبی نرزد اهمیت چندانی ندارد، اما در عین حال درک بهتر پایه های ریاضی و منطقی به صورت عام برای استفاده کنندگان منطق مفید است. در اینجا یکی از مهمترین خواص منطق های چندمقداری، یعنی «معنی شناسی» را مورد مطالعه بیشتر قرار داده ایم. «معنی شناسی» از آن خواصی است که بیشتر مورد توجه محققینی که فقط به کاربردهای منطق های چندمقداری علاقه مند هستند، قرار می گیرد. در اینجا مروری خواهیم داشت بر معنی شناسی منطق های چندمقداری و سپس دوگان معنی شناسی متداول را معرفی خواهیم کرد. همانطور که خواهید دید، این دوگان منجر به پیدایش یک متریک بین گزاره های منطقی خواهد شد که می تواند منجر به پیدایش کاربردهای جدیدی از منطق های چندمقداری شود. به علاوه شاید برای استفاده از ابزارهای آنالیز در منطق نیز مفید واقع گردد.

۲. معنی شناسی متداول منطق های چندمقداری گزاره ای

در معنی شناسی استاندارد منطق های چندمقداری، بازه یکه بسته $[0, 1]$ به عنوان مجموعه مقادیر درستی، عدد صفر به عنوان غلط محض، عدد یک به عنوان درست محض و سایر اعداد به عنوان درست یا غلط نسبی در نظر گرفته می شوند. معمولاً این معنی شناسی استاندارد را منطق فازی می نامند. البته با توجه به نوع کاربردی که از یک منطق فازی مورد انتظار است، روابط منطقی متفاوتی ممکن است در آن منطق مورد استفاده قرار گیرند. می دانیم در منطق کلاسیک همه روابط منطقی عبارتند از:

$$\{\wedge, \vee, \neg, \rightarrow, \leftrightarrow, \perp, \top\}$$

و در عین حال می دانیم به لحاظ تکمیل بودن، این مجموعه روابط منطقی کامل هستند و هر رابط منطقی دیگری توسط آنها قابل تعبیر است و از روی آنها ساخته می شود. حتی به طور دقیق تر می دانیم مجموعه روابط منطقی $\{\wedge, \neg\}$ در منطق کلاسیک گزاره ای کامل بوده و سایر روابط منطقی را می توان از روی آنها تعریف کرد [۱]. اما در منطق های فازی از آنجایی که تعبیر یک رابط منطقی n - موضعی تابعی است از مجموعه $[0, 1]^n$ به $[0, 1]$ ، لذا نه تنها مجموعه روابط منطقی کاملی شبیه منطق کلاسیک قابل معرفی نیست، بلکه طیف روابط منطقی مورد استفاده بسیار متنوع می تواند باشد. یکی از انواع منطق های فازی که بیشتر مورد توجه محققین واقع شده است، منطق های فازی مبتنی بر نرم مثلثی^۴ [۲، ۳] هستند که در آنها روابط منطقی مورد استفاده عبارتند از:

$$\{\wedge, \vee, \&, \oplus, \neg, \rightarrow, \leftrightarrow, \perp, \top\}$$

^۴triangular norm

البته در همین منطق‌ها نیز روابط منطقی فراوان دیگری ممکن است به مجموعه روابط منطقی افزوده شود که از مشهورترین آنها می‌توان به نفی خودیاب^۵، دلتای باز^۶، و ثابت‌های منطقی $\{r : \bar{r} \in (0, 1)\}$ اشاره کرد. دو رابط منطقی $\&$ و $\oplus$ ، معمولاً به عنوان تعبیری برای «و قوی» و «یا قوی» مورد استفاده قرار می‌گیرند. مفهومی شبیه کامل بودن مجموعه روابط منطقی، ولی به شکل خیلی ضعیف‌تر، در بعضی از منطق‌های فازی برقرار است. مثلاً وقتی نرم مثلی مورد استفاده پیوسته باشد، مجموعه روابط منطقی $\{\&, \rightarrow, \perp\}$ به این معنی که قادر به ساختن همه روابط منطقی $\{\&, \vee, \oplus, \neg, \rightarrow, \leftrightarrow, \perp, \top\}$ هستند، کامل می‌باشد. در واقع تعاریف زیر این کامل بودن را ایجاب می‌کنند:

$$\begin{aligned}\varphi \wedge \psi &:= \varphi \& (\varphi \rightarrow \psi) \\ \varphi \vee \psi &:= ((\varphi \rightarrow \psi) \rightarrow \psi) \wedge ((\psi \rightarrow \varphi) \rightarrow \varphi) \\ \neg \varphi &:= \varphi \rightarrow \perp \\ \top &:= \neg \perp \\ \varphi \leftrightarrow \psi &:= (\varphi \rightarrow \psi) \& (\psi \rightarrow \varphi) \\ \varphi \oplus \psi &:= \neg(\neg \varphi \& \neg \psi)\end{aligned}$$

البته از آنجایی که ممکن است «و قوی» خودتوان نباشد، نمادگذاری $\varphi \& \dots \& \varphi^n := \varphi$ نیز ممکن است مورد استفاده قرار گیرد. در سال‌های اخیر اکثر تحقیقات انجام شده در حوزه منطق‌های چندمقداری بر روی منطق‌های مبتنی بر نرم مثلی پیوسته و توسعه‌های آنها بوده است. ما نیز توجه خود را روی این نوع از منطق‌های چندمقداری معطوف می‌کنیم.

تعریف ۱.۲. اگر I یک مجموعه مرتب و $\{p_i\}_{i \in I}$ مجموعه نمادهای گزاره‌ای اتمیک باشد، آنگاه مجموعه همه نمادهای گزاره‌ای تولید شده از روی P توسط روابط منطقی دو موضعی $\&$ و $\rightarrow$ و رابط منطقی صفر موضعی $\perp$ را با P نمایش می‌دهیم. هر عضو P را یک گزاره و هر زیرمجموعه آن را یک تتوری می‌نامیم.

تابع درستی^۹ رابط منطقی «و قوی»، یعنی $\&$ ، یک نرم مثلی پیوسته و تابع درستی $\rightarrow$ نیز عملگر الحاقی آن نرم مثلی پیوسته در نظر گرفته می‌شوند. عملکرد توابع درستی سایر روابط منطقی بر اساس توابع درستی مربوط به روابط منطقی $\&$ و $\rightarrow$ مشخص می‌شود.

تعریف ۲.۲. یک نرم مثلی پیوسته یا t -نرم پیوسته تابعی است مثل $[0, 1] \rightarrow [0, 1]^2 : *$ بطوریکه به ازای هر $x, y, z \in [0, 1]$

$$\begin{aligned}(T1) \quad & x * y = y * x \\ (T2) \quad & (x * y) * z = x * (y * z) \\ (T3) \quad & \text{اگر } x \leq y \text{ آنگاه } x * z \leq y * z \\ (T4) \quad & 1 * x = x\end{aligned}$$

^۵ involutive negation

^۶ نقض نقیض یک گزاره تحت نفی خودیاب، با خود آن گزاره هم معنی است

^۷ Baaz Delta operator

^۸ شبیه رابط منطقی $\top$ که همواره ارزش ۱ دارد، ثابت منطقی $\bar{r}$ ارزش درستی برابر r دارد

^۹ truth function

براحتی می‌توان دید برای هر $x \in [0, 1]$ ، $0 * x = 0$.

مثال ۳.۲. مشهورترین t -نرم‌های پیوسته عبارتند از t -نرم لوکاسویچ، گودل^{۱۰} و حاصل ضربی که به صورت زیر تعریف می‌شوند:

$$\begin{aligned} \bullet \text{ لوکاسویچ: } x *_L y &= \max\{0, x + y - 1\} \\ \bullet \text{ گودل: } x *_G y &= \min\{x, y\} \\ \bullet \text{ حاصل ضربی: } x *_\pi y &= x \cdot y \end{aligned}$$

تعریف ۴.۲. جفت الحاقی^{۱۱} t -نرم $*$ تابعی است مثل $[0, 1]^2 \rightarrow [0, 1]$: $\rightarrow$ بطوریکه به ازای هر $x, y, z \in [0, 1]$ ، $x * z \leq y$ اگر و فقط اگر $y \rightarrow x \rightarrow z$.

توجه کنید که براحتی با توجه به خاصیت کوچکترین کران بالایی اعداد حقیقی و همچنین پیوسته بودن تابع $*$ می‌توان نشان داد $\{z : x * z \leq y\} = \{z : x \rightarrow y\}$.

مثال ۵.۲. عملگر الحاقی t -نرم‌های لوکاسویچ، گودل و حاصل ضربی به صورت زیر می‌باشد:

$$\begin{aligned} \bullet \text{ لوکاسویچ: } x \rightarrow_L y &= \begin{cases} 1 & x \leq y \\ 1 + y - x & x > y \end{cases} \\ \bullet \text{ گودل: } x \rightarrow_G y &= \begin{cases} 1 & x \leq y \\ y & x > y \end{cases} \\ \bullet \text{ حاصل ضربی: } x \rightarrow_\pi y &= \begin{cases} 1 & x \leq y \\ \frac{y}{x} & x > y \end{cases} \end{aligned}$$

توجه کنید که $\rightarrow_L$ پیوسته است در حالی که $\rightarrow_G$ و $\rightarrow_\pi$ پیوسته نمی‌باشند. اکنون می‌توان معنی‌شناسی استاندارد^{۱۲} منطق گزاره‌ای مبتنی بر t -نرم پیوسته را به صورت زیر بیان کرد.

تعریف ۶.۲. فرض کنید $[0, 1] \rightarrow \mathcal{P}$: e . یک تابع باشد.

• می‌توان e را بطور منحصر بفردی به یک تابع $[0, 1] \rightarrow \mathcal{P}$ به صورت زیر تعمیم داد:

$$\begin{aligned} e(\perp) &= 0 - \\ e(\varphi \rightarrow \psi) &:= e(\varphi) \rightarrow e(\psi) - \\ e(\varphi \&\psi) &:= e(\varphi) * e(\psi) - \end{aligned}$$

که آن را یک تابع ارزش^{۱۳} (یا یک ارزیابی) می‌نامند.

• اگر $T \subseteq \mathcal{P}$ و به ازای هر $\varphi \in T$ داشته باشیم $e(\varphi) = 1$ ، می‌نویسیم $e \models T$ و e را مدلی از T می‌نامیم. اگر تئوری T یک مدل داشته باشد، آن را ارضا پذیر^{۱۴} می‌نامند.

^{۱۰}Gödel

^{۱۱}adjoint pair

^{۱۲}standard semantic

^{۱۳}evaluation

^{۱۴}satisfiable

براحتی می‌توان دید $e(\varphi \vee \psi) = \max\{e(\varphi), e(\psi)\}$ و $e(\varphi \wedge \psi) = \min\{e(\varphi), e(\psi)\}$ اما در مورد سایر روابط منطقی، با توجه به t -نرم $*$ ، تابع درستی روابط منطقی ممکن است با هم فرق داشته باشد.

مثال ۷.۲. در منطق لوکاسویچ، چون
$$x \rightarrow_L y = \begin{cases} 1 & x \leq y \\ 1 + y - x & x > y \end{cases}$$
 می‌توان دید $e(\varphi \leftrightarrow \psi) = 1 - |e(\varphi) - e(\psi)|$.

از آنجایی که در منطق لوکاسویچ $e(\neg\varphi) = 1 - e(\varphi)$ ، لذا

$$e(\neg(\varphi \leftrightarrow \psi)) = |e(\varphi) - e(\psi)|$$

به عبارت دیگر تابع درستی $\neg(\varphi \leftrightarrow \psi)$ یک فاصله (متر) بین گزاره‌ها معرفی می‌کند. این در حالی است که در مورد منطق گودل و یا حاصل ضربی تابع درستی $\neg(\varphi \leftrightarrow \psi)$ فقط منجر به متر گسسته می‌شود. همانطور که در بخش آخر مقاله خواهیم دید، معنی‌شناسی دوگان، این مشکل را در منطق گودل و منطق حاصل ضربی حل می‌کند.

۳. منطق فازی مرتبه اول و رابطه تشابه

در حالت منطق فازی مرتبه اول، شبیه منطق مرتبه اول کلاسیک، یک زبان مرتبه اول عبارت است از یک مجموعه به صورت

$$\mathcal{L} = \{ \{ (f_i, n_{f_i}) \}_{i \in I}, \{ (P_j, n_{P_j}) \}_{j \in J} \}$$

که در آن به ازای هر $i \in I$ یک نماد تابعی n_{f_i} -متغیره $(n_{f_i} \in \omega)$ و به ازای هر $j \in J$ یک نماد محمولی n_{P_j} -متغیره $(n_{P_j} \in \omega)$ می‌باشد. با در نظر گرفتن یک مجموعه شمارا از متغیرها، استفاده از روابط منطقی $\{ \&, \rightarrow, \perp \}$ ، سورهای $\{ \forall, \exists \}$ ، نمادهای زبان $\mathcal{L}$ و نماد پرانتز، همه $\mathcal{L}$ -ترم‌ها و $\mathcal{L}$ -فرمول‌ها شبیه منطق کلاسیک ساخته می‌شوند.

مشابه منطق گزاره‌ای، معنی‌شناسی استاندارد بر اساس یک t -نرم پیوسته و عملگر الحاقی آن t -نرم و با در نظر گرفتن بازه یکه بسته $[0, 1]$ به عنوان مجموعه مقادیر درستی صورت می‌گیرد.

تعریف ۸.۳. فرض کنید $\mathcal{L}$ یک زبان مرتبه اول باشد. یک $\mathcal{L}$ -ساختار (یا $\mathcal{L}$ -مدل) $\mathcal{M}$ ، مجموعه‌ای است مثل M (که جهان ساختار نامیده می‌شود) همراه با:

(۱) به ازای هر نماد تابعی $f \in \mathcal{L}$ ، تابعی مثل $f^{\mathcal{M}} : M^{n_f} \rightarrow M$. البته وقتی

$n_f = 0$ ، $f^{\mathcal{M}}$ را عنصری از M در نظر می‌گیریم. معمولاً وقتی $n_f = 0$ ، f را یک

نماد ثابت می‌نامند.

(۲) به ازای هر نماد محمولی $P \in \mathcal{L}$ ، تابعی مثل $P^{\mathcal{M}} : M^{n_P} \rightarrow [0, 1]$. در اینجا

وقتی $n_P = 0$ ، $P^{\mathcal{M}}$ را عنصری از $[0, 1]$ در نظر می‌گیریم.

تعبیر $\mathcal{L}$ -فرمول $\varphi(\bar{x})$ در $\mathcal{L}$ -ساختار $\mathcal{M}$ تابعی است مثل $\varphi^{\mathcal{M}} : M^n \rightarrow [0, 1]$ که شبیه منطق کلاسیک به صورت استقرایی تعریف می‌شود. در اینجا تعبیر سورها به صورت زیر است:

- اگر $\varphi(\bar{x}) = \forall y \psi(y, \bar{x})$ ، آنگاه $\varphi^{\mathcal{M}}(\bar{a}) = \inf_{b \in M} \{ \psi^{\mathcal{M}}(b, \bar{a}) \}$
- اگر $\varphi(\bar{x}) = \exists y \psi(y, \bar{x})$ ، آنگاه $\varphi^{\mathcal{M}}(\bar{a}) = \sup_{b \in M} \{ \psi^{\mathcal{M}}(b, \bar{a}) \}$

اگر ساختاری مثل $\mathcal{M}$ موجود باشد که به ازای هر گزاره φ در تئوری مرتبه اول T ، $\varphi^{\mathcal{M}} = 1$ ، آنگاه $\mathcal{M}$ را مدلی از T می‌نامند و می‌نویسند $\mathcal{M} \models T$.

محمول دوموضوعی تساوی در منطق کلاسیک و نظریه مدل منطق کلاسیک اهمیت ویژه‌ای دارد. خواص اصلی محمول تساوی توسط اصول موضوعه زیر در منطق کلاسیک بیان می‌شوند.

$$\begin{aligned} & \forall x (x = x) \\ & \forall x, \forall y ((x = y) \rightarrow (y = x)) \\ & \forall x, \forall y, \forall z (((x = y) \wedge (y = z)) \rightarrow (x = z)) \end{aligned}$$

محمول تشابه^{۱۵} در واقع نقش محمول تساوی در منطق‌های چندارزشی را ایفا می‌کند. فرض کنید e یک محمول دوموضوعی متعلق به زبان $\mathcal{L}$ باشد، که قرار است نقش محمول تساوی را در منطق فازی داشته باشد. اصول تشابه^{۱۶} در مورد محمول e عبارتند از:

$$\begin{aligned} & \forall x e(x, x) \quad (S1) \\ & \forall x, \forall y (e(x, y) \rightarrow e(y, x)) \quad (S2) \\ & \forall x, \forall y, \forall z ((e(x, y) \wedge e(y, z)) \rightarrow e(x, z)) \quad (S3) \end{aligned}$$

که به اختصار آنها را با C_S نمایش می‌دهیم. قضیه زیر نشان می‌دهد در منطق لوکاسویچ مرتبه اول از روی محمول تشابه e می‌توان یک شبه‌متر بدست آورد.

قضیه ۲.۳. [۳] فرض کنید زبان مرتبه اول $\mathcal{L}$ شامل یک محمول تشابه e باشد. اگر $C_{L\vee}$ نشان دهنده همه اصول منطق لوکاسویچ مرتبه اول باشد و $\mathcal{M} \models C_{L\vee} \cup C_S$ ، آنگاه تعبیر $\neg e$ در $\mathcal{M}$ یک شبه‌متر خواهد بود.

در مورد منطق گودل و منطق حاصلضربی قضیه فوق درست نیست. اما در اینجا نیز دوگان معنی‌شناسی حاضر باعث خواهد شد که تعبیر محمول تشابه e در هر سه نوع منطق چندارزشی لوکاسویچ، گودل و حاصلضربی یک شبه‌متر باشد. شاید به همین دلیل بهتر است این معنی‌شناسی را، معنی‌شناسی متریک^{۱۷} بنامیم و در عوض، معنی‌شناسی که تا الان به کار می‌بردیم را معنی‌شناسی فازی^{۱۸} بنامیم.

۴. نتایج اصلی، معنی‌شناسی متریک منطق‌های چندارزشی

اساس معنی‌شناسی متریک در حالت استاندارد، مبتنی بر در نظر گرفتن $\cdot$ به عنوان درست محض و در مقابل ۱ به عنوان غلط محض می‌باشد. لذا توجه کنید که مثلاً موقع تعریف معنی‌شناسی متریک برای $\varphi \wedge \psi$ ، چون $(\varphi \wedge \psi)^{\mathcal{M}} = \min\{\varphi^{\mathcal{M}}, \psi^{\mathcal{M}}\}$ وقتی درست است^{۱۹} که هر دوی $\varphi^{\mathcal{M}}$ و $\psi^{\mathcal{M}}$ درست باشند، لذا خواهیم داشت:

$$(\varphi \wedge \psi)^{\mathcal{M}} = \max\{\varphi^{\mathcal{M}}, \psi^{\mathcal{M}}\} \quad (20)$$

^{۱۵} similarity relation

^{۱۶} axioms of similarity

^{۱۷} metric semantic

^{۱۸} fuzzy semantic

^{۱۹} برابر صفر است

^{۲۰} این معنی دقیقاً دوگان معنی‌شناسی فازی، یعنی $(\varphi \wedge \psi)^{\mathcal{M}} = \min\{\varphi^{\mathcal{M}}, \psi^{\mathcal{M}}\}$ است

تابع درستی رابط منطقی $\&$ در معنی‌شناسی متریک یک هم‌نرم مثلثی پیوسته^{۲۱} می‌باشد که مثل نرم مثلثی پیوسته تعریف می‌شود با این تفاوت که بجای (T۴) باید داشته باشیم:

$$\cdot *^m x = x$$

هم‌نرم مثلثی پیوسته و جفت الحاقی آن برای منطق‌های لوکاسویچ، گودل و حاصل‌ضربی به صورت زیر می‌باشد:

$x \dot{\rightarrow}_L^m y = \begin{cases} \cdot & x \geq y \\ y - x & x < y \end{cases}$	$x *_L^m y = \min\{1, x + y\}$	لوکاسویچ
$x \dot{\rightarrow}_G^m y = \begin{cases} \cdot & x \geq y \\ y & x < y \end{cases}$	$x *_G^m y = \max\{x, y\}$	گودل
$x \dot{\rightarrow}_\pi^m y = \begin{cases} \cdot & x \geq y \\ \frac{y-x}{x} & x < y \end{cases}$	$x *_\pi^m y = x + y - x.y$	حاصل‌ضربی

معنی‌شناسی متریک منطق مرتبه اول فازی مبتنی بر هم‌نرم مثلثی پیوسته $*^m$ و جفت الحاقی آن، یعنی $\dot{\rightarrow}^m$ ، به صورت زیر است.

تعریف ۱.۴. به ازای هر n تایی $\bar{x}$ ، تعبیر $\mathcal{L}$ -فرمول $\varphi(\bar{x})$ در $\mathcal{L}$ -ساختار $\mathcal{M}$ ، تابعی است مثل $\varphi^{\mathcal{M}}: M^n \rightarrow [0, 1]$ که به صورت استقرایی زیر تعریف می‌شود.

$$(1) \quad \perp^{\mathcal{M}} = 1$$

$$(2) \quad \text{به ازای هر نماد محمولی } n\text{-متغیره } P,$$

$$P^{\mathcal{M}}(t_1(\bar{x}), \dots, t_n(\bar{x}))(\bar{a}) = P^{\mathcal{M}}(t_1^{\mathcal{M}}(\bar{a}), \dots, t_n^{\mathcal{M}}(\bar{a}))$$

$$(3) \quad (\varphi \& \psi)^{\mathcal{M}}(\bar{a}) = \varphi^{\mathcal{M}}(\bar{a}) *_\psi^{\mathcal{M}}(\bar{a})$$

$$(4) \quad (\varphi \rightarrow \psi)^{\mathcal{M}}(\bar{a}) = \varphi^{\mathcal{M}}(\bar{a}) \dot{\rightarrow}^m \psi^{\mathcal{M}}(\bar{a})$$

$$(5) \quad \varphi^{\mathcal{M}}(\bar{a}) = \sup_{b \in M} \{\psi^{\mathcal{M}}(b, \bar{a})\} \text{ آنگاه } \varphi(\bar{x}) = \forall y \psi(y, \bar{x})$$

$$(6) \quad \varphi^{\mathcal{M}}(\bar{a}) = \inf_{b \in M} \{\psi^{\mathcal{M}}(b, \bar{a})\} \text{ آنگاه } \varphi(\bar{x}) = \exists y \psi(y, \bar{x})$$

قضیه زیر برای هر سه نوع منطق لوکاسویچ، گودل و حاصل‌ضربی درست است. برای کمتر شدن متن مقاله آن را برای دو حالت لوکاسویچ و گودل ثابت می‌کنیم.

قضیه ۲.۴. فرض کنید $\mathcal{C}_{L\vee}$ و $\mathcal{C}_{G\vee}$ به ترتیب اصول موضوعه منطق‌های لوکاسویچ و گودل باشند. همچنین فرض کنید $\varphi(\bar{x})$ و $\psi(\bar{x})$ دو فرمول n -متغیره و $\bar{a} \in M^n$.

• اگر در حالت معنی‌شناسی متریک $\mathcal{M} \models \mathcal{C}_{L\vee} \cup \mathcal{C}_S$ ، آنگاه تعبیر e در $\mathcal{M}$ یک شبه متر است و بعلاوه $(\varphi \leftrightarrow \psi)^{\mathcal{M}}(\bar{a}) = |\varphi^{\mathcal{M}}(\bar{a}) - \psi^{\mathcal{M}}(\bar{a})|$.

• اگر در حالت معنی‌شناسی متریک $\mathcal{M} \models \mathcal{C}_{G\vee} \cup \mathcal{C}_S$ ، آنگاه تعبیر e در $\mathcal{M}$ یک شبه فرامتر^{۲۲} است و بعلاوه $(\varphi \leftrightarrow \psi)^{\mathcal{M}}(\bar{a}) = d_{\max}(\varphi^{\mathcal{M}}(\bar{a}), \psi^{\mathcal{M}}(\bar{a}))$ که در آن،

$$d_{\max}(x, y) = \begin{cases} \max\{x, y\} & x \neq y \\ \cdot & x = y \end{cases}$$

^{۲۱} continuous t -conorm

^{۲۲} شبه متریک که رابطه مثلث بصورت قوی $d(a, b) \leq \max\{d(a, c), d(b, c)\}$ در آن برقرار است

برهان. از اینکه $\mathcal{M} \models \forall x e(x, x)$ داریم $\sup_{a \in M} e^{\mathcal{M}}(a, a) = \cdot$. لذا برای هر $a \in M$ نتیجه می‌گیریم $\hat{\cdot} = e^{\mathcal{M}}(a, a)$. از طرف دیگر چون $\mathcal{M} \models \forall x, \forall y (e(x, y) \rightarrow e(y, x))$ لذا برای هر $a, b \in M$ داریم $e^{\mathcal{M}}(a, b) \geq e^{\mathcal{M}}(b, a)$ که با عوض کردن جای a و b نتیجه می‌گیریم $e^{\mathcal{M}}(a, b) = e^{\mathcal{M}}(b, a)$. نامساوی مثلث و ادامه اثبات را در هر حالت جداگانه بررسی می‌کنیم.

حالت ۱: اگر $\mathcal{M} \models \mathcal{C}_{L\vee} \cup \mathcal{C}_S$ ، آنگاه از (S^3) داریم

$$\cdot \sup_{a,b,c \in M} [(e^{\mathcal{M}}(a, c) *_L^m e^{\mathcal{M}}(c, b)) \dot{\rightarrow}_L^m e^{\mathcal{M}}(a, b)] = \cdot$$

لذا برای هر $a, b, c \in M$ خواهیم داشت $e^{\mathcal{M}}(a, b) \leq e^{\mathcal{M}}(a, c) + e^{\mathcal{M}}(c, b)$ که نشان می‌دهد $e^{\mathcal{M}}$ یک شبه متر است. از طرفی چون

$$(\varphi \rightarrow \psi)^{\mathcal{M}}(\bar{a}) = \begin{cases} \cdot & \varphi^{\mathcal{M}}(\bar{a}) \geq \psi^{\mathcal{M}}(\bar{a}) \\ \psi^{\mathcal{M}}(\bar{a}) - \varphi^{\mathcal{M}}(\bar{a}) & \varphi^{\mathcal{M}}(\bar{a}) < \psi^{\mathcal{M}}(\bar{a}) \end{cases}$$

براحتی می‌توان دید

$$\begin{aligned} (\varphi \leftrightarrow \psi)^{\mathcal{M}}(\bar{a}) &= ((\varphi \rightarrow \psi) \wedge (\psi \rightarrow \varphi))^{\mathcal{M}}(\bar{a}) \\ &= \max \{(\varphi \rightarrow \psi)^{\mathcal{M}}, (\psi \rightarrow \varphi)^{\mathcal{M}}\} \\ &= |\varphi^{\mathcal{M}}(\bar{a}) - \psi^{\mathcal{M}}(\bar{a})| \end{aligned}$$

حالت ۲: چنانچه $\mathcal{M} \models \mathcal{C}_{G\vee} \cup \mathcal{C}_S$ ، آنگاه $\mathcal{M} \models (S^3)$ ایجاب می‌کند که

$$\cdot \sup_{a,b,c \in M} [(e^{\mathcal{M}}(a, c) *_G^m e^{\mathcal{M}}(c, b)) \dot{\rightarrow}_G^m e^{\mathcal{M}}(a, b)] = \cdot$$

پس برای هر $a, b, c \in M$ خواهیم داشت $e^{\mathcal{M}}(a, b) \leq \max\{e^{\mathcal{M}}(a, c), e^{\mathcal{M}}(c, b)\}$ که نشان می‌دهد $e^{\mathcal{M}}$ یک شبه فرامتر است. اتحاد دوم نیز براحتی از روی تعریف $\dot{\rightarrow}_G^m$ نتیجه می‌شود. $\square$

قضیه بالا نشان می‌دهد تعبیر رابط منطقی $\leftrightarrow$ در منطق‌های لوکاسویچ و گودل (و حاصل ضربی) یک متریک می‌باشد. با استفاده از شبه متر و متر مذکور در صورت قضیه قبل و روش ابرضرب، ما یکی از مهمترین قضیه‌های معنی‌شناسی، یعنی قضیه فشردگی را برای منطق‌های مذکور اثبات کردیم. این ایده برای بسط نظریه مدل روی منطق‌های مذکور نیز می‌تواند به کار گرفته شود. در منطق پیوسته مرتبه اول که تعمیمی از منطق لوکاسویچ و با معنی‌شناسی متریک می‌باشد، این کار انجام شده است، اما برای منطق گودل، منطق حاصل ضربی و به طور کلی منطق‌های مبتنی بر نرم مثلثی پیوسته در حال مطالعه می‌باشد.

مراجع

۱. م. اردشیر، منطق ریاضی، انتشارات هرمس، ویرایش دوم، ۱۳۹۱.
۲. ا. اسلامی، منطق فازی و کاربردهای آن، انتشارات دانشگاه شهید باهنر کرمان، ۱۳۹۱.
3. Hajek, P. (1998). *Metamathematics of fuzzy logic*, Springer Science.



معرفی شاخص‌های کارایی فازی چندمتغیره

بنیتا دولت‌زاده^۱ *، مصطفی گریوانی^۲، و ماشالله ماشین‌چی^۳

^۱دانشگاه شهید باهنر کرمان

^۲دانشگاه شهید باهنر کرمان

mostafa.gerivani@ymail.com

^۳دانشگاه شهید باهنر کرمان

چکیده. ما در این مقاله قصد داریم به معرفی شاخص‌های کارایی چندمتغیره برای زمانی که حدود مشخصه‌ی فنی فازی هستند، بپردازیم. در ابتدا مرور کلی بر چند شاخص رایج کارایی فرایند چندمتغیره داریم. سپس تعاریفی را در زمینه علم فازی ارائه می‌دهیم و در نهایت با بکارگیری علم فازی و استفاده از شاخص‌های کارایی چندمتغیره، به معرفی شاخص‌های کارایی فازی چندمتغیره می‌پردازیم.

۱. مقدمه

شاخص‌های کارایی چندمتغیره جهت ارزیابی وضعیت فرایندهای تولیدی با چندین مشخصه همبسته (نظیر وزن، طول و عرض) مورد استفاده می‌باشند. با گسترش سیستم‌های اتوماتیک کنترل محصول و امکان دسترسی به حجم بالاتری از داده‌های کمی و کیفی در رابطه با فرایند، شاخص‌های کارایی چندمتغیره کاربرد وسیع‌تری نسبت به شاخص‌های کارایی تک‌متغیره در صنعت دارند (شهریاری و عبدالله‌زاده، ۱۳۸۹). از طرفی امروزه بدلیل وجود داده‌ها و اطلاعات فازی در فرایندهای در حال تولید، استفاده از این علم برای محاسبه و بیان شاخص‌های کارایی واژگان کلیدی. شاخص‌های کارایی چندمتغیره، حدود مشخصه‌ی فنی فازی، شاخص‌های کارایی فازی

چندمتغیره .

* سخنران

فرایند چندمتغیره امری مهم قلمداد می‌شود. با در نظر گرفتن حدود مشخصه فنی فازی، می‌توان از رخ دادن خطاهای احتمالی جلوگیری کرد. به همین دلیل در این مقاله از ترکیب علم فازی و شاخص‌های کارایی چندمتغیره برای زمانی که حدود مشخصه فنی فازی باشند، استفاده کرده و شاخص‌های چندمتغیره‌ی جدیدی را در محیط فازی معرفی می‌کنیم.

۲. رایج‌ترین شاخص‌های کارایی چندمتغیره

تعریف ۱.۲. شاخص C_{pm} بسط شاخص کارایی فرایند دومتغیره مطرح شده توسط هیوبل و همکاران (۱۹۹۱) است، که توسط شهریار و همکارانش (۱۹۹۵) به صورت زیر مطرح شده است (شهریار و عبدالله‌زاده، ۱۳۸۹)

$$C_{pm} = \left[\frac{\prod_{i=1}^v (USL_i - LSL_i)}{\prod_{i=1}^v (UPL_i - LPL_i)} \right]^{\frac{1}{v}} \quad (1.2)$$

در رابطه فوق USL_i و LSL_i به ترتیب حد بالا و پایین مؤلفه‌ی اصلی i ام می‌باشند. همچنین هاردل و سیمار (۲۰۰۷) برای محاسبه‌ی UPL_i و LPL_i دو فرمول زیر را مطرح کردند

$$UPL_i = \mu_i + \sqrt{\chi_{\alpha,p}^2 \cdot \sigma_i^2}, \quad LPL_i = \mu_i - \sqrt{\chi_{\alpha,p}^2 \cdot \sigma_i^2}$$

که در آنها μ_i و σ_i^2 به ترتیب میانگین و واریانس مؤلفه‌ی اصلی i ام و $\chi_{\alpha,p}^2$ چندک مرتبه α ام از توزیع کای-دو با p درجه آزادی است (هاردل و سیمار، ۲۰۰۷).

شاخص C_{pm} عملاً برابر است با حجم ناحیه تیرانس به حجم ناحیه فرایند تعدیل شده و چون این نسبت به اندازه‌ی v بستگی دارد، از ریشه v ام این نسبت استفاده شده است. لازم به ذکر است این شاخص اولین مؤلفه‌ی بردار کارایی فرایند چندمتغیره است که این بردار را به صورت زیر نمایش می‌دهند

$$MPCV = [C_{pm}, PV, LI]$$

مؤلفه‌ی دوم جهت بررسی موقعیت مرکز فرایند نسبت به مقدار هدف تعریف گشته است. این مؤلفه مقدار $P - value$ آزمون $H_0: \mu = T$ در مقابل $H_1: \mu \neq T$ می‌باشد که در آن $\mu = (\mu_1, \dots, \mu_v)'$ و $T = (T_1, \dots, T_v)'$ است. در این حالت با سطح معنی‌داری $\alpha = 0.05$ ، هرگاه $PV > \alpha$ باشد بیانگر این است که مرکز فرایند با مقدار هدف بر هم منطبق هستند. این امر نشان‌دهنده‌ی کارایی مطلوب فرایند از نظر نزدیکی به مقدار هدف در سطح معنی‌داری مورد نظر می‌باشد. مؤلفه‌ی سوم شاخص مکانی می‌باشد که آن با نماد LI نشان می‌دهند. این شاخص در بردارنده اطلاعاتی در رابطه با موقعیت مکانی ناحیه فرایند تعدیل شده نسبت به ناحیه تفرانس می‌باشد. هنگامی که ناحیه فرایند تعدیل شده کاملاً درون ناحیه تفرانس باشد، $LI = 1$ و در غیر این صورت $LI = 0$ (وانگ و همکاران، ۲۰۰۰).

تعریف ۲.۲. شاخص $MC_{p(pc)}$ با استفاده از تجزیه و تحلیل مؤلفه‌های اصلی توسط وانگ و چن (۱۹۹۹) به صورت زیر مطرح شده است

$$MC_{p(pc)} = \sqrt[k]{\prod_{i=1}^v C_{p;pc_i}}$$

که در آن $PC_i = e'_i X; i = 1, \dots, v$ مؤلفه‌ی اصلی i ام، k برابر با تعداد مؤلفه‌های اصلی مورد نظر در تجزیه و تحلیل و $C_{p;pc_i}$ شاخص کارایی فرایند متناظر با i امین مؤلفه‌ی اصلی می‌باشد و از رابطه زیر قابل محاسبه است:

$$C_{p;pc_i} = \frac{USL_{pc_i} - LSL_{pc_i}}{S_{pc_i}} \quad (2.2)$$

که در این رابطه S_{pc_i} برابر با واریانس نمونه‌ای مؤلفه‌ی اصلی i ام است. همچنین محاسبه حدود مشخصه پایین و بالا، مقدار هدف و واریانس نمونه‌ای مؤلفه‌های اصلی به صورت زیر محاسبه می‌شود

$$LSL_{pc_i} = e'_i \cdot LSL, \quad USL_{pc_i} = e'_i \cdot USL, \quad T_{pc_i} = e'_i \cdot T, \quad S_{pc_i} = e'_i \cdot S_{e_i}$$

که در آن $USL = (USL_1, \dots, USL_v)'$ و $LSL = (LSL_1, \dots, LSL_v)'$ است (وانگ و دیو، ۲۰۰۰).

تعریف ۳.۲. شاخص S_{pk}^T توسط چن و همکاران در سال ۲۰۰۳ ارائه شده است (عبادی و امیری، ۲۰۱۲)

$$S_{pk}^T = \frac{1}{3} \phi^{-1} \left\{ \frac{[\prod_{j=1}^v (2\Phi(3S_{pk_j}) - 1) + 1]}{2} \right\}$$

که در آن Φ تابع توزیع تجمعی نرمال استاندارد می‌باشد. این شاخص بر مبنای احتمال تولید محصول نامنطبق مطرح شده است و بسط شاخص تک‌متغیره S_{pk} می‌باشد که توسط بویلز در سال (۱۹۹۴) به صورت زیر ارائه شده است

$$S_{pk} = \frac{1}{3} \phi^{-1} \left\{ \frac{1}{2} \Phi \left(\frac{USL - \mu}{\sigma} \right) + \frac{1}{2} \Phi \left(\frac{\mu - LSL}{\sigma} \right) \right\}$$

۳. شاخص‌های کارایی فازی چندمتغیره

تعریف ۱.۳. فرض کنید $\widetilde{USL} : \mathbb{R} \rightarrow [0, 1]$ یک مجموعه‌ی فازی و $\mathbb{R}$ مجموعه‌ی اعداد حقیقی باشد که در شرایط زیر صدق می‌کند؛
الف) $\widetilde{USL}$ یک تابع غیرصعودی باشد،

ب) وجود داشته باشد $u_1 \in \mathbb{R}$ بطوریکه به ازای هر $x \leq u_1$ داشته باشیم $\widetilde{USL}(x) = 1$ ،
آنگاه $\widetilde{USL}$ را یک کران بالای فازی می‌نامند (پرچی و ماشین‌چی، ۱۳۸۹) و (پرچی و همکاران، ۲۰۱۰).

تعریف ۲.۳. فرض کنید $\widetilde{LSL} : \mathbb{R} \rightarrow [0, 1]$ یک مجموعه‌ی فازی و $\mathbb{R}$ مجموعه‌ی اعداد حقیقی باشد که در شرایط زیر صدق می‌کند؛
الف) $\widetilde{LSL}$ یک تابع غیرنزولی باشد،

ب) وجود داشته باشد $l_1 \in \mathbb{R}$ بطوریکه به ازای هر $x \geq l_1$ داشته باشیم $\widetilde{LSL}(x) = 1$ ،

آنگاه $\widetilde{LSL}$ را یک کران پایین فازی می‌نامند (پرچمی و ماشین‌چی، ۱۳۸۹) و (پرچمی و همکاران، ۲۰۱۰).

تعریف ۳.۳. تفاضل میان دو کران فازی $\widetilde{USL}$ و $\widetilde{LSL}$ را با نماد $\widetilde{USL} \ominus \widetilde{LSL}$ نمایش می‌دهیم و آن را به صورت زیر تعریف می‌کنیم

$$\widetilde{USL} \ominus \widetilde{LSL} = \int_0^1 g(\alpha) (u_\alpha - l_\alpha) d\alpha \quad (۱.۳)$$

که در آن به ازای هر α در بازه $[0, 1]$ داریم

$$\widetilde{USL}_\alpha = (-\infty, u_\alpha] \quad , \quad \widetilde{LSL}_\alpha = [l_\alpha, +\infty)$$

و u_0 و l_0 اعداد حقیقی متناهی و g تابعی غیرنزولی روی بازه $[0, 1]$ باشد بطوریکه $g(0) = 0$

$$\text{و} \quad \int_0^1 g(\alpha) d\alpha = 1$$

(پرچمی و ماشین‌چی، ۱۳۸۹).

با توجه به تعاریف ذکر شده درباره حدود مشخصه‌ی فنی فازی و بیان تفاضل میان این دو کران فازی و همچنین با توجه به شاخص‌های کارایی فرایند چندمتغیره، می‌توان شاخص‌های کارایی فازی چندمتغیره را به صورت زیر معرفی کرد.

تعریف ۴.۳. شاخص $C_{\widetilde{pm}}$ را می‌توان با توجه به تعریف ۱.۲ و رابطه (۱.۳) به صورت زیر تعریف کرد

$$C_{\widetilde{pm}} = \left[\frac{\prod_{i=1}^v (\widetilde{USL}_i \ominus \widetilde{LSL}_i)}{\prod_{i=1}^v (UPL_i - LPL_i)} \right]^{\frac{1}{v}}.$$

تعریف ۵.۳. شاخص $MC_{\widetilde{p(pc)}}$ را می‌توان با توجه به تعریف ۲.۲ و رابطه (۱.۳) به صورت

زیر تعریف کرد

$$MC_{\widetilde{p(pc)}} = \sqrt[k]{\prod_{i=1}^k C_{\widetilde{p;pc}_i}}$$

که در آن $PC_i = e'_i X; i = 1, \dots, v$ مؤلفه‌ی اصلی i ام، k برابر با تعداد مؤلفه‌های اصلی مورد نظر در تجزیه و تحلیل و $C_{\widetilde{p};pc_i}$ شاخص کارایی فرایند متناظر با i امین مؤلفه‌ی اصلی می‌باشد که با توجه به رابطه (۱.۳) به صورت زیر تعریف می‌شود

$$C_{\widetilde{p};pc_i} = \frac{\widetilde{USL}_{pc_i} \ominus \widetilde{LSL}_{pc_i}}{S_{pc_i}}.$$

تعریف ۶.۳. شاخص S_{pk}^T را می‌توان با توجه به تعاریف ۳.۲، ۱.۳ و ۲.۳ به صورت زیر تعریف نمود

$$S_{pk}^T = \frac{1}{3} \phi^{-1} \left\{ \frac{[\prod_{j=1}^v (2\Phi(3\widetilde{S}_{pk_j}) - 1) + 1]}{2} \right\}$$

که در آن

$$\widetilde{S}_{pk} = \frac{1}{3} \phi^{-1} \left\{ \frac{1}{2} \Phi \left(\frac{\widetilde{USL} \ominus \mu}{\sigma} \right) + \frac{1}{2} \Phi \left(\frac{\mu \ominus \widetilde{LSL}}{\sigma} \right) \right\}.$$

در خصوص بیان رابطه‌ی بین شاخص‌های کارایی معمولی چند متغیره و شاخص‌های کارایی فازی چند متغیره، که در این مقاله پیشنهاد شده‌است، می‌توان به این مطلب اشاره کرد که برای ساخت شاخص‌های کارایی فازی چند متغیره، از شاخص‌های کارایی معمولی چند متغیره به‌گونه‌ای که تفاضل فازی حدود فنی را برای آن در نظر بگیریم استفاده شده‌است.

لازم به تذکر است که شاخص‌های کارایی چند متغیره به ویژه زمانی که حدود مشخصه‌ی فنی فازی باشد، در صنایع مختلف همچون صنعت پزشکی، کشاورزی و ... کاربرد دارد. بطورکلی کیفیت محصولات و تولیداتی که دارای شرایط ذکر شده در مقاله هستند، با این شاخص‌ها بررسی می‌شوند.

بحث و نتیجه‌گیری

ما در این مقاله به معرفی شاخص‌های کارایی فازی چندمتغیره برای زمانی که حدود مشخصه‌ی فنی فازی باشند، پرداختیم. بنابراین با در نظر گرفتن حدود مشخصه‌ی فنی فازی، می‌توان از رخ

دادن خطاهای احتمالی جلوگیری کرد. همچنین برای بررسی فرایندهای تولیدی با چندین مشخصه همبسته، استفاده از شاخص‌های کارایی چندمتغیره کیفیت محصول را دقیق‌تر کنترل می‌کند. به همین دلیل در این مقاله شاخص‌های کارایی فازی چندمتغیره را معرفی کردیم تا در صورت فازی بودن حدود مشخصه‌ی فنی بتوان از این شاخص‌ها استفاده کرد.

مراجع

- پرچمی، ع. و ماشین‌چی، م. (۱۳۸۹). کنترل کیفیت آماری، دانشگاه شهید باهنر کرمان.
- شهریاری، ح. و عبدالله‌زاده، م. (۱۳۸۹). معرفی شاخص جدید کارایی فرایند چند متغیره، چهارمین کنفرانس بین‌المللی مهندسی صنایع، دانشگاه تربیت مدرس .
- Ebadi M. and Amiri A. (2012), Evaluation of process capability in multi-variate simple linear profiles, *Sharif University of Technology Scientia Iranica*, 19(6), 1960 - 1968.
- Hardle W. and Simar L. (2007), Applied multivariate statistical analysis, *Springer Verlag*, Berlin.
- Parchami A. and Mashinchi M. and Sharayei A. (2010), An effective approach for measuring the capability of manufacturing processes, *Production Planning and Control*, 21(3), 250 - 257.
- Wang F.k. and Du T.C.T. (2000), Using principal component analysis in process performance for multivariate data, *The International Journal of Management Science*, omega 28.

Wang F.k. and Hubele N.F. and Lawrence F.P and Miskulin J.D and Shahri-
ari H. (2000), Comparison of three multivariate process capability in-
dices, *Journal of Quality Technology*, 32(3).

رگرسیون ریج فازی در حضور داده‌های پرت برای ضرایب و خروجی فازی

محمدرضا ربیعی^{۱*}، معصومه فرخی^۲، و محمد آرشی^۲

^۱شاهرود، دانشگاه صنعتی شاهرود، دانشکده علوم ریاضی، گروه آمار
rabiei1354@yahoo.com

^۲شاهرود، دانشگاه صنعتی شاهرود، دانشکده علوم ریاضی، گروه آمار

چکیده. در این مقاله یک روش رگرسیون ریج فازی جدید برای زمانی که متغیرهای ورودی دقیق و خروجی فازی باشد، پیشنهاد شده است. سپس با معرفی معیارهای نیکویی برازش مدل پیشنهادی را با مدل‌های دیگر مورد ارزیابی قرار می‌دهیم و سپس با ارائه یک مثال عددی برتری مدل پیشنهادی را زمانی که در میان داده‌ها داده‌ی پرت نیز وجود دارد، نشان می‌دهیم.

۱. مقدمه

تحلیل رگرسیون یک روش آماری برای بررسی و مدل‌سازی رابطه بین متغیرهاست. این روش تقریباً در کلیه رشته‌های علوم از جمله مهندسی، فیزیک، اقتصاد، مدیریت، علوم زیستی، کشاورزی و علوم اجتماعی مورد استفاده واقع می‌شود. درحقیقت، تحلیل رگرسیون یکی از کاربردی‌ترین روش‌های آماری است. مدل‌های رگرسیونی در حوزه‌های کاربردی مورد استفاده قرار می‌گیرند. یک مسئله جدی که می‌تواند استفاده از مدل رگرسیونی را با مشکل مواجه کند، همخطی چندگانه یا وابستگی خطی نزدیک بین متغیرهای رگرسیونی است. وجود وابستگی خطی نزدیک، توانایی برآورد ضرایب مدل رگرسیون را با مشکل مواجه می‌کند. روش‌های متعددی برای غلبه بر مشکل همخطی در مدل‌های خطی ارائه شده است. برخی روش‌های اریب عبارتند از: برآوردیابی انقباضی، مولفه‌های اصلی، رویکرد ریج و یکی از موثرترین روش‌ها برای حل مشکل همخطی، رگرسیون ریج است که توسط هورل و کنارد [۹] در سال ۱۹۷۰ معرفی شد.

یکی دیگر از مفروضات در رگرسیون کلاسیک این است که متغیرهای مورد مطالعه، متغیرهایی دقیق هستند و مشاهدات مربوط به متغیرها نیز مشاهداتی دقیق باشند. زمانی که این فرض برقرار نباشد رگرسیون فازی جایگزین رگرسیون کلاسیک می‌شود. رگرسیون فازی اولین بار توسط تاناکا [۱] در سال ۱۹۸۰ معرفی شد. در زمینه رگرسیون فازی مطالعات زیادی صورت گرفته است که از جمله آن‌ها می‌توان به موارد زیر اشاره کرد: ایشیوشی و تاناکا [۱۱] مدل‌های رگرسیون فازی غیرخطی را با توجه به کاربرد آن در شبکه‌های عصبی در سال ۱۹۹۲ پیشنهاد داده‌اند. یک

واژگان کلیدی. رگرسیون فازی، ضریب تعیین تعمیم یافته، همخطی چندگانه، رگرسیون ریج فازی.
* سخنران

روش مهم دیگر برای حل دستگاه رگرسیون فازی، روش رگرسیون حداقل کمترین توان‌های دوم است که توسط دیاموند [۴] پیشنهاد شده است.

در رگرسیون فازی نیز مانند رگرسیون کلاسیک زمانی که بین متغیرهای پیش‌بین مشکل هم‌خطی وجود داشته باشد، روش رگرسیون ریج فازی یک روش موثر برای رفع این مشکل است. اولین بار هانگ و همکاران [۱۰] الگوریتم یادگیری رگرسیون ریج که توسط ساندروز و همکاران [۱۲] برای متغیرهای دوگانه مورد بررسی قرار گرفته بود را برای مدل‌های رگرسیون فازی با مقادیر ورودی دقیق و خروجی اعداد فازی نرمال تعمیم دادند. همچنین دونسو و همکاران [۵] رگرسیون ریج فازی را برای مدل‌های درجه دوم پیشنهاد داده‌اند. بلاسوندارام و کاپیل [۲] نیز یک روش رگرسیون ریج فازی وزن‌دار شده برای متغیرهای ورودی دقیق و خروجی فازی پیشنهاد داده‌اند. ساختار مقاله به صورت زیر است: در بخش ۲ مفاهیم و تعاریف اولیه فازی را به طور مختصر شرح می‌دهیم، در بخش ۳ مدل رگرسیون ریج فازی را پیشنهاد می‌دهیم، در بخش ۴ چندین معیار نیکویی برازش را بیان می‌کنیم و در بخش ۵ با ارائه دو مثال مدل رگرسیون ریج فازی پیشنهادی را مورد بررسی قرار می‌دهیم.

۲. مفاهیم و تعاریف اولیه

در این بخش به ارائه تعاریف و قضایای مورد نیاز در زمینه اعداد فازی و خواص بین آن‌ها می‌پردازیم.

تعریف ۱.۲. هر عدد فازی مثلثی $A = (m_A, l_A, r_A)_T$ با تابع عضویت $A(\cdot)$ به صورت زیر تعریف می‌شود:

$$A(x) = \begin{cases} \frac{x - (m_A - l_A)}{l_A} & m_A - l_A \leq x < m_A \\ \frac{(m_A + r_A) - x}{r_A} & m_A \leq x < m_A + r_A \\ 0 & \text{otherwise} \end{cases}$$

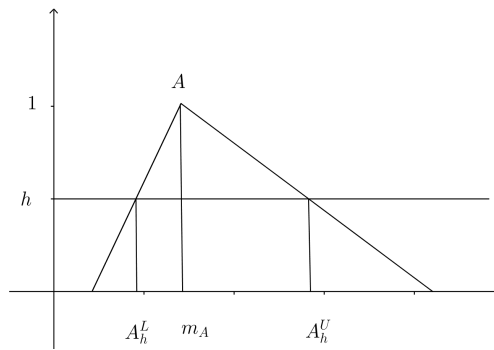
که m_A مرکز، l_A و r_A به ترتیب پهنای چپ و پهنای راست عدد فازی است. مجموعه تمام اعداد فازی مثلثی را با $T(\mathbb{R})$ نمایش می‌دهند.

تعریف ۲.۲. برای هر عدد فازی $A = (m_A, l_A, r_A)_T \in T(\mathbb{R})$ و $h \in [0, 1]$ ، h -برش A به صورت زیر تعریف می‌شود:

$$[A^L(h), A^U(h)] = [m_A - (1 - h)l_A, m_A + (1 - h)r_A] \quad (۱.۲)$$

قضیه ۳.۲. [۷] اگر $A = (m_A, l_A, r_A)_T$ و $B = (m_B, l_B, r_B)_T$ و $\lambda \in \mathbb{R}$ آنگاه بر اساس اصل گسترش، داریم:

$$\begin{aligned} A + B &= (m_A + m_B, l_A + l_B, r_A + r_B)_T \\ A - B &= (m_A - m_B, l_A + r_B, r_A + l_B)_T \\ \lambda \cdot A &= \begin{cases} (\lambda m_A, \lambda l_A, \lambda r_A)_T & \lambda > 0 \\ (\lambda m_A, -\lambda r_A, -\lambda l_A)_T & \lambda < 0 \end{cases} \end{aligned} \quad (۲.۲)$$



شکل ۱: h -برش عدد فازی A

۱.۲. فاصله بین دو عدد فازی. دو عدد مثلثی $B = (m_B, l_B, r_B)_T$ و $A = (m_A, l_A, r_A)_T$ ترسیم شده در شکل ۲ را در نظر بگیرید. فاصله بین این دو عدد فازی به شکل های مختلف محاسبه می شود که به چند روش مرسوم زیر اشاره می شود.

۱.۱.۲. فاصله وزن دار شده. ژو^۱ [۱۴] فاصله ی وزن دار شده ی بین دو عدد فازی بر پایه ی تابع وزنی $f(h) = h$ را به صورت زیر تعریف می کنند:

$$\begin{aligned} d_1^2(A, B) = & (m_A - m_B)^2 \\ & + \frac{1}{3} (m_A - m_B) [(r_A - r_B) - (l_A - l_B)] \\ & + \frac{1}{12} [(l_A - l_B)^2 + (r_A - r_B)^2] \end{aligned} \quad (3.2)$$

۲.۱.۲. فاصله دونسو. دونسو^۲ و همکاران [۶] فاصله ی بین دو عدد فازی را به صورت زیر تعریف می کنند:

$$\begin{aligned} d_2^2(A, B) = & k_1 (m_A - m_B)^2 \\ & + k_2 ((m_A - l_A - (m_B - l_B))^2 + (m_A + r_A - (m_B + r_B))^2) \end{aligned} \quad (4.2)$$

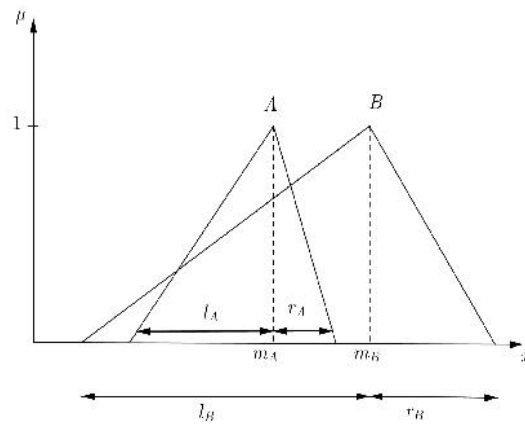
که در آن k_1 و k_2 اعداد حقیقی بین صفر و یک هستند که توسط محقق تعیین می شوند. در صورتی که $k_1 > k_2$ اهمیت مرکز بیشتر از پهنای پاسخ ها در نظر گرفته می شود و اگر $k_1 < k_2$ پهنای پاسخ ها از اهمیت بیشتری برخوردار می گردد. همچنین در صورتی که $k_1 = k_2$ ، اهمیت مرکز و پهنای یکسان فرض شده است.

۳.۱.۲. فاصله حسن پور. حسن پور و همکاران [۸] فاصله بین دو عدد فازی را به صورت زیر بیان کردند:

$$d_3(A, B) = |m_A - m_B| + |l_A - l_B| + |r_A - r_B| \quad (5.2)$$

^۱Xu

^۲Donoso



شکل ۲: دو عدد فازی مثلی و فاصله بین آنها

گزاره ۴.۲. [۸] تابع d_3 یک متر روی $T(\mathbb{R}) \times T(\mathbb{R})$ است. یعنی برای A, B, C داریم:

1. $d_3(A, B) \geq 0, d_3(A, A) = 0$
2. $d_3(A, B) = d_3(B, A)$
3. $d_3(A, C) \leq d_3(A, B) + d_3(B, C)$

توجه ۵.۲. برای دو عدد دقیق فاصله‌های تعریف شده فوق به قدرمطلق تفاوت بین آنها خلاصه می‌شود.

۳. مدل رگرسیون ریج فازی پیشنهادی

مدل رگرسیون خطی فازی زیر را در نظر بگیرید:

$$Y_i = A_0 + A_1 x_{i1} + \dots + A_p x_{ip}, \quad i = 1, \dots, n \quad (1.3)$$

که در آن x_{ij} ها مقادیر دقیق و $A_j = (m_{A_j}, l_{A_j}, r_{A_j})_T$ به ازای $j = 1, 2, \dots, p$ ضرایب رگرسیونی فازی و $Y_i = (m_{Y_i}, l_{Y_i}, r_{Y_i})_T$ به ازای $i = 1, 2, \dots, n$ متغیر پاسخ فازی هستند.

بدون کاستن از کلیت مسأله فرض می‌کنیم $x_{ij} > 0$ (در صورتی که x_{ij} ها اعدادی منفی باشند می‌توان با یک انتقال آن‌ها را به اعدادی مثبت تبدیل کرد). با استفاده از قضیه ۳.۲ مدل (۱.۳) به صورت زیر بدست می‌آید

$$\begin{aligned} \hat{Y}_i &= (m_{A_0}, l_{A_0}, r_{A_0})_T + (m_{A_1}, l_{A_1}, r_{A_1})_T x_{i1} + \dots + (m_{A_p}, l_{A_p}, r_{A_p})_T x_{ip} \\ &= (m_{A_0} + \sum_{j=1}^p m_{A_j} x_{ij}, l_{A_0} + \sum_{j=1}^p l_{A_j} x_{ij}, r_{A_0} + \sum_{j=1}^p r_{A_j} x_{ij})_T \end{aligned} \quad (2.3)$$

بنابراین با توجه به فاصله (۵.۲) مسئله بهینه سازی مدل رگرسیون ریج فازی را برپایه منطق مدل رگرسیون ریج در آمار به صورت زیر می توان نوشت:

$$\begin{aligned} & \min_{A_0, A_1, \dots, A_p} \sum_{i=1}^n d(Y_i, \hat{Y}_i) + \lambda \sum_{j=0}^p A_j^2 \\ & = \min_{A_0, A_1, \dots, A_p} \sum_{i=1}^n (|m_{Y_i} - m_{\hat{Y}_i}| + |l_{Y_i} - l_{\hat{Y}_i}| + |r_{Y_i} - r_{\hat{Y}_i}|) + \lambda \sum_{j=0}^p A_j^2 \end{aligned} \quad (3.3)$$

۴. معیارهای نیکویی برازش

در رگرسیون فازی یک موضوع مهم، توان تشریح مدل می باشد. در این بخش، تعدادی معیار نیکویی برازش برای اندازه گیری برازش یک مدل رگرسیونی معرفی می کنیم.

۱.۴. ضریب تعیین تعمیم یافته. یکی از شاخص های ارزیابی نیکویی برازش مدل رگرسیون آماری، محاسبه ضریب تعیین می باشد. با الهام از این تعریف در آمار کلاسیک، برای مدل های رگرسیونی فازی، ضریب تعیین تعمیم یافته را می توان به صورت زیر تعریف کرد.

تعریف ۱.۴. در مدل رگرسیون خطی فازی اگر Y_i و $\hat{Y}_i$ به ترتیب متغیر وابسته و مقدار پیش بینی آن توسط مدل رگرسیون فازی باشد، ضریب تعیین تعمیم یافته برای مدل را به صورت زیر تعریف می کنیم

$$R_G^2 = \frac{\sum_{i=1}^n d^2(\hat{Y}_i, \bar{Y})}{\sum_{i=1}^n d^2(Y_i, \bar{Y})} \quad (1.4)$$

۲.۴. میانگین خطای پیشگویی. میانگین خطای پیشگویی^۳ (MPE) به عنوان معیار دیگری برای ارزیابی نیکویی برازش استفاده کرد. در زیر معیار MPE را در قالب یک تعریف می آوریم.

تعریف ۲.۴. تحت مفروضات مدل رگرسیونی فازی (۱.۳) اگر $\hat{Y}_i$ مقدار پیشگویی Y_i باشد آنگاه میانگین خطای پیشگویی (MPE) برابر است با

$$MPE = \frac{1}{n} \sum_{i=1}^n d(Y_i, \hat{Y}_i) \quad (2.4)$$

۳.۴. نمودار حبابی. کپی و همکاران [۳] از نمودار حبابی برای ارزیابی مدل رگرسیون برازش داده شده با خروجی فازی، استفاده کرده اند. در این نمودار محور x و محور y به ترتیب مراکز مقادیر مشاهده شده و برآورد شده را نشان می دهد. این دایره ها دارای مرکزی به صورت $(m_{Y_i}, m_{\hat{Y}_i})$ و قطری به اندازه $|l_{Y_i} - l_{\hat{Y}_i}| + |r_{Y_i} - r_{\hat{Y}_i}|$ هستند. هر چه مراکز این دایره ها به نیمساز ربع اول و سوم نزدیک تر باشند مقادیر برآورد شده به مقادیر مشاهده شده به خوبی برازش داده شده است و هر چه این دایره ها کوچکتر باشند، پهنای عدد فازی برآوردهای دقیقتری هستند.

^۳Mean prediction error

۵. مثال کاربردی

در این بخش رفتار مدل رگرسیون ریح فازی پیشنهادی را در یک مثال کاربردی مورد بررسی قرار می‌دهیم.

داده‌های سیمان پرتلند. در این مثال داده‌های وود^۴ و همکاران [۱۳] را مورد استفاده قرار می‌دهیم (جدول ۱). این داده‌ها از مطالعه تأثیر ترکیبات مختلف روی حرارت ساطع شده در طول سخت شدن سیمان پرتلند بدست آمده است. در اینجا داده‌هایی که از چهار ترکیب $3CaO.Al_2O_3$ ، $3CaO.SiO_2$ ، (x_2) ، $4CaO.Al_2O_3Fe_2O_3$ ، (x_3) ، $2CaO.SiO_2$ ، (x_4) و میزان گرمای ساطع شده بعد از 80° روز (y) را مورد بررسی قرار می‌دهیم. با توجه به برخی ابهامات در اندازه‌گیری‌های مربوط به مشاهدات متغیر پاسخ، با مشورت یک فرد متخصص با استفاده از یک عدد فازی مثلثی با پهنای راست و چپ متناسب با y فازی سازی شده‌اند (جدول ۱).

جدول ۱: داده‌های سیمان پرتلند

ردیف	x_1	x_2	x_3	x_4	Y	$\hat{Y}$	d_3
1	7	26	6	60	$(78.5, 8.63, 10.20)_T$	$(78.50, 8.63, 13.51)_T$	3.31
2	1	29	15	52	$(74.3, 8.91, 14.86)_T$	$(72.44, 11.13, 10.41)_T$	8.51
3	11	56	8	20	$(104.3, 19.81, 11.47)_T$	$(104.71, 13.95, 14.80)_T$	9.60
4	11	31	8	47	$(87.6, 13.14, 14.89)_T$	$(88.62, 9.87, 14.89)_T$	4.29
5	7	52	6	33	$(95.9, 15.34, 16.30)_T$	$(95.74, 12.89, 13.52)_T$	5.37
6	11	55	9	22	$(109.2, 12.01, 16.38)_T$	$(105.33, 14.04, 14.93)_T$	7.34
7	3	71	17	6	$(102.7, 12.32, 11.29)_T$	$(105.24, 18.48, 11.29)_T$	8.71
8	1	31	22	44	$(72.5, 13.05, 10.15)_T$	$(76.80, 13.05, 10.15)_T$	4.30
9	2	54	18	22	$(93.1, 14.89, 14.89)_T$	$(91.82, 15.90, 10.67)_T$	6.50
10	21	47	4	26	$(115.911.59, 20.86)_T$	$(115.9, 11.59, 19.35)_T$	1.51
11	1	40	23	34	$(83.3, 15.82, 9.99)_T$	$(83.29, 14.75, 10.13)_T$	1.21
12	11	66	9	12	$(113.3, 15.86, 13.59)_T$	$(113.3, 15.86, 15.06)_T$	1.47
13	10	68	8	12	$(109.4, 18.59, 12.03)_T$	$(112.56, 15.96, 14.71)_T$	8.47
MPE							5.434

چون ورودی‌ها دقیق هستند از VIF معمولی برای نشان دادن میزان همخطی میان متغیرهای مستقل استفاده می‌کنیم که

$$VIF_1 = 38, VIF_2 = 254, VIF_3 = 46, VIF_4 = 250$$

با توجه به بزرگتر از 10° بودن همه مقادیر VIF ، وجود همخطی میان متغیرهای مستقل تایید می‌شود و بنابراین لزوم استفاده از رگرسیون ریح فازی محرز است. برای برآورد مقادیر پاسخ با استفاده از نرم افزار متمتیکا نسخه 9.0° نیاز به انتخاب پارامتر ریح λ داریم. با توجه به ازای

^۴Wood

مقادیر مختلف λ کمترین مقدار MPE به ازای $\lambda = 0.2$ است، بنابراین $\lambda = 0.2$ به عنوان پارامتر ریح انتخاب می‌کنیم. ضرایب مدل بصورت زیر بدست می‌آید:

$$A_0 = (-0.059, 1.900, 1.895)_T$$

$$A_1 = (2.205, 0.009, 0.500)_T$$

$$A_2 = (1.151, 0.176, 0.094)_T$$

$$A_3 = (0.831, 0.236, 0.038)_T$$

$$A_4 = (0.470, 0.011, 0.091)_T$$

در جدول ۲ برای مدل‌های معرفی شده شاخص MPE محاسبه شده است که MPE_1 برای داده‌های اولیه است و MPE_2 برای زمانی است که داده‌ی دهم به مقدار ۳۰ واحد افزایش یافته و به داده پرت تبدیل شده است، همچنین MPE_3 برای زمانی است که علاوه بر داده‌ی دهم، داده‌ی ششم نیز به مقدار ۳۰ واحد افزایش یافته است و به داده پرت تبدیل شده‌اند. نتایج جدول نشان دهنده آن است که در هر صورت مدل پیشنهادی برازش بهتری را به داده‌ها در مقابل مدل‌های قبلی نشان می‌دهد.

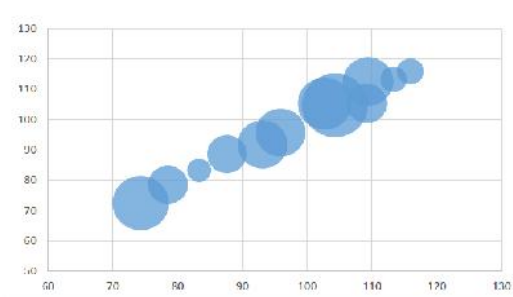
جدول ۲: مقایسه MPE

مدل	MPE_1	MPE_2	MPE_3
d_1 بر اساس متر	5.809	9.748	11.067
d_2 بر اساس متر	5.776	8.868	11.063
d_3 بر اساس متر	5.434	7.909	10.362

در شکل ۳ نمودار حبابی آمده است که در آن محور افقی، مقادیر Y_i و محور عمودی، مقادیر $\hat{Y}_i$ آمده‌است. همانطور که پیشتر بیان شد هر چه مراکز این دایره‌ها به نیمساز ربع اول و سوم نزدیک‌تر باشند مقادیر برآورد شده به مقادیر مشاهده شده به خوبی برازش داده شده است که این مساله در شکل کاملاً مشهود است. همچنین هر چه این دایره‌ها کوچکتر باشند، پهنای اعداد فازی برآوردهای دقیقتری هستند. همچنین مقدار شاخص ضریب تعیین تعمیم یافته معرفی شده در (۱۰۴) برابر 0.901 است که میزان بسیار خوبی است.

۶. نتیجه‌گیری

در این مقاله یک روش رگرسیون ریح فازی برای زمانی که داده‌های ورودی دقیق و داده‌های خروجی فازی مثلی باشند، پیشنهاد داده‌ایم. نتایج حاکی از آن است که اگرچه زمانی که داده پرت در بین داده‌ها وجود ندارد مدل پیشنهادی برپایه متر d_3 کمی بهتر از دو مدل دیگر است اما با افزایش داده‌های پرت کارایی مدل کاملاً مشهود است و مدل پیشنهادی از استواری خوبی نسبت به مدل‌های موجود برخوردار است.



شکل ۳: نمودار حبابی

REFERENCES

1. H Tanaka-S Uegima-K Asai, *Linear regression analysis with fuzzy model*, IEEE Trans. Systems Man Cybern **12** (1982), 903–907.
2. S Balasundaram and Kapil, *Weighted fuzzy ridge regression analysis with crisp inputs and triangular fuzzy outputs*, International Journal of Advanced Intelligence Paradigms **3** (2011), no. 1, 67–81.
3. Renato Coppi, Pierpaolo D’Urso, Paolo Giordani, and Adriana Santoro, *Least squares estimation of a linear regression model with lr fuzzy response*, Computational Statistics & Data Analysis **51** (2006), no. 1, 267–286.
4. Phil Diamond, *Fuzzy least squares*, Information Sciences **46** (1988), no. 3, 141–157.
5. Sergio Donoso, Nicolás Marín, and MA Vila, *Fuzzy ridge regression with non symmetric membership functions and quadratic models*, International Conference on Intelligent Data Engineering and Automated Learning, Springer, 2007, pp. 135–144.
6. ———, *Fuzzy ridge regression with non symmetric membership functions and quadratic models*, International Conference on Intelligent Data Engineering and Automated Learning, Springer, 2007, pp. 135–144.
7. Didier J Dubois, *Fuzzy sets and systems: theory and applications*, vol. 144, Academic press, 1980.
8. H Hassanpour, HR Maleki, and MA Yaghoobi, *Fuzzy linear regression model with crisp coefficients: a goal programming approach*, Iranian Journal of Fuzzy Systems **7** (2010), no. 2, 1–153.
9. Arthur E Hoerl and Robert W Kennard, *Ridge regression: Biased estimation for nonorthogonal problems*, Technometrics **12** (1970), no. 1, 55–67.
10. Dug Hun Hong, Changha Hwang, and Chulhwan Ahn, *Ridge estimation for regression models with crisp inputs and gaussian fuzzy output*, Fuzzy Sets and Systems **142** (2004), no. 2, 307–319.
11. Hisao Ishibuchi and Hideo Tanaka, *Fuzzy regression analysis using neural networks*, Fuzzy sets and systems **50** (1992), no. 3, 257–265.
12. Craig Saunders, Alexander Gammerman, and Volodya Vovk, *Ridge regression learning algorithm in dual variables.*, ICML, vol. 98, 1998, pp. 515–521.

13. Hubert Woods, Harold H Steinour, and Howard R Starke, *Effect of composition of portland cement on heat evolved during hardening*, Industrial & Engineering Chemistry **24** (1932), no. 11, 1207–1214.
14. R Xu, *A linear regression model in fuzzy environment*, Advance in Modelling Simulation **27** (1991), no. 2.



تحلیل رگرسیون استوار فازی براساس کمترین قدرمطلق خطاها و رتبه‌بندی مجموعه‌های فازی

محمد رضا ربیعی^۱، الهام رحیمیان^{۱*}، و داود شاهسونی^۱

^۱شاهرود، دانشگاه صنعتی شاهرود، دانشکده علوم ریاضی، گروه آمار

چکیده. زمانی که مشاهدات فازی باشند و در مجموعه داده‌ها، مشاهدات دورافتاده وجود داشته باشند از رگرسیون جایگزین استوار فازی برای مدل‌بندی استفاده می‌کنیم. در این مقاله تحلیل رگرسیون کمترین قدر مطلق فازی اصلاح شده‌ای را، برای حالتی که ضرایب رگرسیونی اعداد فازی و مشاهدات متغیرهای مستقل، دقیق باشند و در مجموعه داده‌ها مشاهدات دور افتاده داشته باشیم، مطرح می‌کنیم. در این روش برای مقایسه مجموعه‌های فازی، باقی‌مانده‌ها رتبه‌بندی می‌شوند. سپس ماتریس وزن توسط تابع عضویت باقیمانده‌ها تعریف می‌شود و برآوردهای کمترین قدرمطلق فازی موزون با استفاده از ماتریس وزن بدست می‌آیند. برای نشان دادن عملکرد روش پیشنهادی، مثالی مطرح و نتایج و مقایسات حاصل ارائه می‌شود.

۱. مقدمه

روش کمترین توان‌های دوم (LS)، روشی معمول است که سال‌هاست برای برآورد پارامترهای رگرسیون به‌کار می‌رود. اگر مفروضات روش LS برای یک مجموعه داده برقرار باشد، برآوردهای روش LS، به‌عنوان بهترین برآوردها شناخته می‌شوند. ولی زمانی که در مجموعه داده‌ها، مشاهدات دورافتاده وجود داشته باشند و متغیرها یا مشاهدات مربوط به آن‌ها نادقیق باشند از تحلیل رگرسیون جایگزین استوار فازی استفاده می‌شود. اولین پژوهش در مورد تحلیل رگرسیون خطی فازی توسط تاناکا و همکاران [۱۲] انجام شد. دیاموند [۴] روش کمترین توان‌های دوم فازی (FLS)^۲ را پیشنهاد کرد. چانگ و لی [۳] برای مواردی که در آن‌ها داده‌های دورافتاده حضور دارند، تعمیمی از روش کمترین توان‌های دوم موزون فازی را پیشنهاد کردند که به هر درجه عضویت وزنی را اختصاص می‌دهد و بر پایه‌ی یک اثر متقابل با نظر تصمیم‌گیرنده است. اما همه این روش‌ها به داده‌های دورافتاده حساس هستند و نیاز است که از روش‌های استوار استفاده کنیم. در تحلیل رگرسیون فازی بر مبنای کمترین قدر مطلق انحرافات (LAD)^۳ می‌توان به مقالات [۹، ۱۰، ۱۱، ۱۳] اشاره کرد. رحیمیان و همکاران [۱] مطالعه‌ای تحت عنوان «تحلیل رگرسیون استوار فازی با داده‌های خروجی و پارامترهای فازی بر پایه رتبه‌بندی مجموعه‌های

واژگان کلیدی. رگرسیون استوار، داده دور افتاده، رگرسیون فازی، رتبه‌بندی مجموعه‌های فازی .
* سخنران

^۱Least- Squares

^۲Fuzzy Least Squares

^۳Least Absolute Deviation

فازی» انجام دادند. آن‌ها برای حالتی که متغیرهای وابسته و ضرایب رگرسیونی اعداد فازی بوده و مجموعه داده‌ها حاوی مشاهدات دور افتاده است، تحلیل رگرسیون کمترین توان‌های دوم فازی اصلاح شده‌ای را مطرح کردند. در این روش با استفاده از یک الگوریتم تکراری چند مرحله‌ای، مقادیر اولیه پارامترهای مدل رگرسیون از روش کمترین توان‌ای دوم با متر دیاوند بدست آمده و برای مقایسه مجموعه‌های فازی، باقی‌مانده‌ها با استفاده از شاخص حضور سراسری برای هر مجموعه فازی (4OM) به دست آمده و رتبه‌بندی شدند. سپس ماتریس وزن توسط تابع عضویت باقیمانده‌ها به شکل دوزنقه‌ای تعریف شده و برآوردهای کمترین توان‌های دوم فازی موزون با استفاده از ماتریس وزن بدست آمده، محاسبه شدند.

در این مقاله با تعمیم روش رحیمیان و همکاران [۱] با استفاده از روش کمترین قدرمطلق خطاها با متر حسن پور مقادیر اولیه محاسبه و با معرفی تابع عضویت جدید برای باقیمانده‌ها مدل رگرسیونی مناسبتری را در مواجهه با داده‌های پرت ارائه می‌دهیم. ساختار کلی این مقاله به این صورت تنظیم شده است که در بخش دوم، به بررسی رگرسیون استوار و M -برآوردگرها پرداخته، در بخش سوم روش رگرسیون فازی مورد بحث قرار می‌گیرد. فاصله بین دو عدد فازی بیان شده و در ادامه رگرسیون کمترین مجموع قدرمطلق خطای فازی معرفی می‌شود. در بخش چهارم رتبه‌بندی مجموعه‌های فازی مورد بررسی قرار می‌گیرد و در بخش بعد مدل رگرسیون استوار فازی پیشنهادی مطرح شده و با مثالی کارایی مدل را نشان خواهیم داد. در بخش پایانی نتیجه گیری را بیان می‌کنیم.

۲. رگرسیون استوار

۱۰۲. روش کمترین توان‌های دوم موزون. یک مدل رگرسیون خطی، بصورت ماتریسی به شکل زیر نمایش داده می‌شود:

$$Y = X\beta + \epsilon$$

$$= \begin{pmatrix} x_{11} & x_{12} & \dots & x_{1p} \\ x_{21} & x_{22} & \dots & x_{2p} \\ \vdots & \vdots & \dots & \vdots \\ x_{n1} & x_{n2} & \dots & x_{np} \end{pmatrix} \begin{bmatrix} \beta_1 \\ \beta_2 \\ \vdots \\ \beta_p \end{bmatrix} + \begin{bmatrix} \epsilon_1 \\ \epsilon_2 \\ \vdots \\ \epsilon_n \end{bmatrix} \quad (10.2)$$

که $Y = (Y_1, Y_2, \dots, Y_n)'$ بردار مشاهدات متغیر وابسته، $X_{n \times p}$ ماتریس مشاهدات متغیرهای مستقل، $\epsilon_{n \times 1}$ بردار خطای تصادفی و $\beta_{p \times 1}$ بردار ضرایب رگرسیونی است. برآورد LS ضرایب رگرسیونی به صورت زیر بدست می‌آید:

$$\hat{\beta} = (X'X)^{-1}X'Y \quad (20.2)$$

روش LS معمول، فرض می‌کند که واریانس‌ها خطای ثابت دارند. روش کمترین توان‌های دوم موزون (WLS)^۵ زمانی مورد استفاده قرار می‌گیرد که فرض ثابت بودن واریانس خطا برقرار نباشد. مدل مورد بررسی (۱۰۲) را با ماتریس قطری واریانس-کوواریانس غیر ثابت در

^۴Overall Measurement

^۵Weighted Least Squares

نظر بگیرید. با فرض $w_i = \frac{1}{\sigma_i^2}, i = 1, 2, \dots, n$ ، عناصر ماتریس قطری W که مبین وزن هاست، برآورد کمترین توان های دوم موزون به صورت زیر به دست می آید:

$$\hat{\beta}_{WLS} = (X'WX)^{-1}X'WY \quad (3.2)$$

۲.۲. M - برآوردگرها. شیوه M - برآوردگرها^۶ مینیم کردن توابع مختلف $\rho(z_i)$ است که در آن $z_i = \frac{e_i}{\hat{\sigma}}$ و $e_i = y_i - x_i'\hat{\beta}$ باقیمانده ی i -امین مشاهده است، $\hat{\beta}$ بردار برآورد ضرایب مدل رگرسیونی و $\hat{\sigma}$ برآورد واریانس خطای مدل است. در حالتی که خطاها به طور نرمال توزیع شده باشند انحرافات استاندارد نمونه، $\hat{\sigma}$ ، اندازه های مناسبی برای تغییرپذیری می باشند، ولی زمانی که در مجموعه داده ها مشاهدات دورافتاده داریم به اندازه های استوار دیگری از تغییرپذیری، نیاز است. یکی از اندازه های استوار تغییرپذیری، به صورت زیر تعریف می شود:

$$\hat{\sigma}_{MAD} = d = \frac{\text{median}|e_i - \text{median}(e_i)|}{0.6745}, \quad i = 1, \dots, n \quad (4.2)$$

بنابراین گوئیم یک برآوردگر استوار است اگر عبارت

$$\sum_{i=1}^n \rho\left(\frac{e_i}{d}\right) \quad (5.2)$$

را مینیم کند. پس، باقیمانده های استاندارد شده جدید، به صورت $z_i = \frac{e_i}{d}$ تعریف می شوند. است. ملاحظه می شود که اگر $\rho(z) = z^2$ باشد با مینیم کردن رابطه (۵.۲) به همان مدل کمترین توان های دوم معمول می رسم (برای مطالعه بیشتر به [۵] مراجعه کنید). روش کمترین توان های دوم در به دست آوردن پارامترها به هر مشاهده، وزن یکسانی می دهد. روش های استوار قادرند به مشاهدات، وزن های نابرابر اختصاص دهند. به طور کلی مشاهداتی که مانده های بزرگی را تولید می کنند به وسیله برآورد استوار، کم وزن ترند. پس می توان گفت که انتخاب $\rho(z)$ خاص، در حقیقت انتخاب تابع وزنی است که می خواهیم برای رتبه بندی مانده ها با اندازه های مختلف تخصیص دهیم. با گرفتن اولین مشتق جزیی از معادله (۵.۲) نسبت به هر $\hat{\beta}_j$ و با برابر صفر قرار دادن آن به شکل زیر، ضرایب رگرسیونی به دست می آیند:

$$\sum_{i=1}^n x_{ij} \Psi \left[\frac{y_i - \sum_{j=1}^p x_{ij} \beta_j}{d} \right] = 0, \quad j = 1, 2, \dots, n \quad (6.2)$$

که $\Psi(z) = d(\rho(z))/d(z)$. توابع زیادی از Ψ ، توسط افراد مختلفی تعریف شده است (برای اطلاعات بیشتر به [۷] مراجعه کنید).

۳. اعداد فازی و رتبه بندی بین آنها

در این بخش به بررسی مختصری از اعداد فازی و فاصله های بین آنها می پردازیم. سپس به تعریفی از رتبه بندی اعداد فازی معرفی شده توسط [۲] می پردازیم و به یک گزاره از آن اشاره می کنیم.

^۶"M" for "Maximum liklihood-type"

۱.۳. حساب اعداد فازی و فاصله بین آن‌ها. در این قسمت به معرفی عدد فازی مثلثی و روابط بین آن‌ها می‌پردازیم.

تعریف ۱.۳. اگر A یک عدد فازی مثلثی با تابع عضویت $A(\cdot)$ باشد به شکل $A = (m_A, l_A, r_A)_T$ نشان داده شده و به صورت زیر تعریف می‌شود:

$$A(x) = \begin{cases} \frac{x - (m_A - l_A)}{l_A} & m_A - l_A \leq x \leq m_A \\ \frac{(m_A + r_A) - x}{r_A} & m_A \leq x \leq m_A + r_A \\ 0 & \text{OW} \end{cases}$$

که در آن $m_A \in \mathbb{R}$ و $l_A, r_A \geq 0$ به ترتیب نشان دهنده مقدار مرکز، پهنای چپ و پهنای راست عدد فازی مثلثی است. مجموعه تمام اعداد فازی مثلثی را با $\mathbb{T}(\mathbb{R})$ نمایش می‌دهند. اگر پهنای چپ و راست با هم برابر باشند یعنی $l_A = r_A = s_A$ ، عدد فازی مثلثی را متقارن گویند که با نماد $(m_A, s_A)_T$ نشان داده می‌شوند.

تعریف ۲.۳. $(h\text{-برش‌ها})$: مجموعه عناصری را از X (مجموعه مرجع) که درجه عضویت آن‌ها در مجموعه فازی A دست کم به بزرگی h ($0 < h \leq 1$) باشد، $h\text{-برش}$ A گوئیم و با A_h نشان می‌دهیم؛ یعنی $A_h = \{x \in X | A(x) \geq h\}$. با توجه به این که حدود پایین و بالای عدد فازی مثلثی A به ترتیب $m_A - l_A$ و $m_A + r_A$ می‌باشند، A_h به صورت زیر به دست می‌آید:

$$[A_L^{-1}(h), A_U^{-1}(h)] = [m_A - (1 - h)l_A, m_A + (1 - h)r_A] \quad (۱.۳)$$

جمع دو عدد فازی و ضرب یک عدد دقیق در یک عدد فازی، بر اساس مفهوم اصل توسیع به دست می‌آید.

گزاره ۳.۳. اگر $A_1 = (m_1, l_1, r_1)_T$ و $A_2 = (m_2, l_2, r_2)_T$ دو عدد فازی در $\mathbb{T}(\mathbb{R})$ باشند و $\lambda \in \mathbb{R}$ آن‌گاه

$$A_1 + A_2 = (m_1 + m_2, l_1 + l_2, r_1 + r_2)_T$$

$$\lambda A_1 = \begin{cases} (\lambda m_1, \lambda l_1, \lambda r_1)_T & \lambda > 0 \\ (\lambda m_1, -\lambda r_1, -\lambda l_1)_T & \lambda < 0 \end{cases}$$

تعریف ۴.۳. (فاصله حسن پور [۸]): اگر داشته باشیم $X = (m_x, l_x, r_x)_T$ و $Y = (m_y, l_y, r_y)_T$ آنگاه فاصله بین دو عدد فازی مثلثی که با $d(X, Y)$ نشان می‌دهیم به صورت زیر تعریف می‌شود:

$$d_1(X, Y) = |m_X - m_Y| + |l_X - l_Y| + |r_X - r_Y| \quad (۲.۳)$$

تعریف ۵.۳. (فاصله وزن دار شده [۱۴]) فرض کنید A و B دو عدد فازی مثلی باشند آنگاه بر پایه تابع وزنی $f(h) = h$ داریم:

$$\begin{aligned} d_2^2(A, B) &= (m_A - m_B)^2 \\ &+ \frac{1}{3}(m_A - m_B)[(r_A - r_B) - (l_A - l_B)] \\ &+ \frac{1}{12}[(l_A - l_B)^2 + (r_A - r_B)^2] \end{aligned} \quad (۳.۳)$$

اگر $A = (m_A, s_A)_T$ و $B = (m_B, s_B)_T$ دو عدد فازی متقارن باشند آنگاه رابطه (۳.۳) به صورت زیر به دست می آید:

$$d_2^2(A, B) = (m_A - m_B)^2 + \frac{1}{6}(s_A - s_B) \quad (۴.۳)$$

۲.۳. رتبه بندی مجموعه های فازی. اعداد فازی مثل اعداد واقعی، نظم و ترتیب خطی ندارند و مقایسه آن ها ساده نیست. بسیاری از محققان در مورد روش های رتبه بندی فازی مطالعاتی انجام داده اند. در این قسمت، روش مقایسه اعداد فازی را که توسط چانگ و لی [۲] مطرح شده است، به اختصار بیان می کنیم. شاخص زیر به منظور مقایسه ی دو عدد فازی مثلی، تعریف می شود.

تعریف ۶.۳. فرض کنید A و B دو عدد فازی مثلی به ترتیب با توابع عضویت $A(x)$ و $B(y)$ باشند. همچنین فرض کنید $A_L^{-1}(h)$ و $A_U^{-1}(h)$ نیز کران های پایین و بالای تصویر تابع عضویت در سطح h بر روی تکیه گاه باشند. تفاضل بین A و B به صورت زیر تعریف می شود:

$$d_3(A, B) = \int_0^1 g_A(\{A^{-1}(h)\}) dh - \int_0^1 g_B(\{B^{-1}(h)\}) dh \quad (۵.۳)$$

که در آن

$$g_A(\{A^{-1}(h)\}) = \omega(h) (\chi_1(h)A_L^{-1}(h) + \chi_2(h)A_U^{-1}(h))$$

$$g_B(\{B^{-1}(h)\}) = \omega(h) (\chi_1(h)B_L^{-1}(h) + \chi_2(h)B_U^{-1}(h))$$

و $A_L^{-1}, A_U^{-1}, B_L^{-1}, B_U^{-1}$ از تعریف ۲.۳ بدست می آید.

توجه ۷.۳. اندازه های وزن $\omega(h)$ و $\chi_1(h)$ و $\chi_2(h)$ می بایست به طور ذهنی توسط تصمیم گیرنده معین شوند. پیشنهادهایی برای $\omega(h)$ به صورت های زیر داده شده است:

$$\omega(h) = \frac{h}{\int_0^1 h dh} \quad \text{یا} \quad \omega(h) = \frac{h(1-h)}{\int_0^1 h(1-h) dh}$$

همچنین اندازه های وزنی ذهنی $\chi_1 = \chi_1(h)$ و $\chi_2 = \chi_2(h)$ به ازای $\chi_1, \chi_2 \in [0, 1]$ را معمولاً در ساده ترین حالت طوری انتخاب می کنند که $\chi_1 + \chi_2 = 1$ باشد (برای اطلاعات بیشتر به [۱] مراجعه کنید).

گزاره ۸.۳. تحت شرایط تعریف ۶.۳ و توجه ۷.۳ اگر $B = (m_B, s_B)_T$ و $A = (m_A, s_A)_T$ اعداد فازی مثلی متقارن و $\chi_1 = \chi_2 = \chi = \frac{1}{2}$ باشند، داریم:

$$E = d_3(A, B) = m_B - m_A \quad (۶.۳)$$

۴. مدل رگرسیون استوار فازی پیشنهادی

در این بخش ابتدا شیوه رگرسیون کمترین مجموع قدرمطلق خطای فازی با متر حسن‌پور را معرفی می‌کنیم. سپس با استفاده از این شیوه در گام اول الگوریتم رحیمیان و همکاران [۱]، به محاسبه ضرایب اولیه مدل پرداخته و با معرفی تابع عضویت جدید برای باقیمانده‌ها در گام چهارم الگوریتم، مدل رگرسیونی بهتری را نسبت به مدل‌های قبلی ارائه می‌دهد.

۱.۴. رگرسیون کمترین مجموع قدر مطلق خطای فازی. هدف آن است که بر پایه‌ی مشاهداتی به صورت $(Y_i, x_{i1}, x_{i2}, \dots, x_{ip})$ که $x_{ij} \in \mathbb{R}$ ($i = 1, \dots, n; j = 1, \dots, p$) و Y_i ها اعداد فازی مثلی به شکل $(m_{Y_i}, l_{Y_i}, r_{Y_i})_T$ هستند، یک مدل بهینه با ضرایب فازی به صورت

$$Y = A_0 + A_1 x_1 + \dots + A_p x_p \quad (۱.۴)$$

که در آن $A_j = (m_{A_j}, l_{A_j}, r_{A_j})_T; j = 0, 1, \dots, p$ ، برای توصیف و تحلیل داده‌ها و پیش‌بینی بر پایه‌ی آن، به دست آوریم.

روش کمترین مجموع قدر مطلق خطای فازی^۷ را مشابه روش کمترین قدر مطلق خطای معمولی، با استفاده از فاصله بین دوعدد فازی تعریف شده در رابطه (۴.۳) و مینیم کردن مجموع فاصله مشاهدات از مقادیر برآورد شده آن‌ها، به دست آمده و پارامترهای مدل برآورد می‌شوند. یعنی

$$M(A_0, A_1, \dots, A_p) = \min \sum_{i=1}^n d_1(A_0 + A_1 x_{i1} + \dots + A_p x_{ip}, Y_i) \quad (۲.۴)$$

با استفاده از گزاره ۳.۳ و با فرض مثبت بودن x_i ها به ازای $i = 1, \dots, n$ (در صورت منفی بودن با اضافه کردن یک مقدار مثبت به داده‌ها، آن‌ها را مثبت می‌کنیم) مدل فوق به شکل زیر به دست می‌آید:

$$(۳.۴)$$

$$y = (m_{A_0} + m_{A_1} x_1 + \dots + m_{A_p} x_p, l_{A_0} + l_{A_1} x_1 + \dots + l_{A_p} x_p + r_{A_0} + r_{A_1} x_1 + \dots + r_{A_p} x_p)_T$$

بنابراین داریم:

$$d_1(A_0 + A_1 x_{i1} + \dots, A_p x_{ip}, Y_i) = |m_{A_0} + m_{A_1} x_{i1} + \dots + m_{A_p} x_{ip} - m_{Y_i}| + |l_{A_0} + l_{A_1} x_{i1} + \dots + l_{A_p} x_{ip} - l_{Y_i}| + |r_{A_0} + r_{A_1} x_{i1} + \dots + r_{A_p} x_{ip} - r_{Y_i}| \quad (۴.۴)$$

می‌توان پارامترها را توسط مشتق جزئی نسبت به m_{A_j} و l_{A_j} و r_{A_j} ها ($j = 0, 1, \dots, p$) و برابر صفر قرار دادن آن‌ها یا به عبارت دیگر به وسیله مینیم سازی رابطه‌ی (۲.۴)، برآورد کرد. در نتیجه FLAD، بر حسب برآورد پارامترها به دست می‌آید.

^۷Fuzzy Least Absolute Deviation (FLAD)

در تحلیل رگرسیون چندگانه، وقتی x_i دقیق و Y_i یک عدد فازی مثلثی متقارن باشد و در مجموعه داده‌ها، مشاهده دور افتاده داشته باشیم، برای مدل‌های رگرسیون برآورد شده، الگوریتم روش رگرسیون استوار فازی به صورت زیر پیشنهاد می‌شود؛ $\hat{A}$ برآورد پارامترهای رگرسیون، k تعداد تکرارها و $\epsilon > 0$ عدد بسیار کوچکی است که توسط تصمیم گیرنده تعیین می‌شود (اینجا مقدار ϵ ، 0.2 در نظر گرفته شد).

الگوریتم: این الگوریتم دارای گام‌های زیر است:

گام ۱: وقتی x_i دقیق و $Y_i = (m_{Y_i}, l_{Y_i}, r_{Y_i})_T$ یک عدد فازی مثلثی باشد، از رابطه‌ی (۲.۴)، برآورد اولیه FLAD را بیابید.

گام ۲: مقادیر برآورد مرکز عدد فازی $(m_{\hat{Y}_i})$ و باقیمانده‌ها (E_i) محاسبه می‌شوند. باقیمانده‌ها را از معادله (۶.۳)، به دست آورید.

گام ۳: میانه بر حسب قدرمطلق مقدار باقیمانده‌ها تعیین می‌شود. فاصله‌ها را از دستور زیر محاسبه کنید:

$$D_i = ||abs(E_i) - median(abs(E_i))||, \quad i = 1, 2, \dots, n \quad (5.4)$$

که $||.||$ فاصله اقلیدسی است.

گام ۴: تابع عضویت جدید را برای باقیمانده‌ها به شکل زیر تعریف می‌کنیم

$$\mu(E) = \begin{cases} 1 & abs(E) \leq a \\ e^{-\left(\frac{abs(E)-a}{d}\right)^2} & abs(E) > a \end{cases} \quad (6.4)$$

که در آن $a = median(D_i)$.

گام ۵: مقادیر عضویت را با استفاده از گام ۴ محاسبه کرده و ماتریس وزن برای هر تابع عضویت که در گام ۴ معرفی شد را تشکیل دهید. این ماتریس‌های وزن، ماتریس‌هایی قطری هستند که عناصر روی قطرهای، درجه عضویت‌ها می‌باشند. پس از آن، برآوردهای WFLAD را از طریق این ماتریس‌های وزن به دست آورید.

گام ۶: اگر $|\hat{A}^{k+1} - \hat{A}^k| < \epsilon$ ، الگوریتم را متوقف کنید. در غیر این صورت به گام ۲ بروید.

توجه ۱.۴. تفاوت الگوریتم این مقاله با مقالات قبلی همانطور که پیشتر گفته شد در گام‌های ۱ و ۴ است. در گام ۴ تابع عضویت باقی مانده‌ها در مطالعات ذکر شده به صورت زیر است:

$$\mu(E) = \begin{cases} 1 & abs(E) \leq a \\ \frac{b-abs(E)}{b-a} & a < abs(E) < b \\ 0 & OW \end{cases} \quad (7.4)$$

که در آن $b = max(D_i) + d$ و d ، پارامتر استوار تغییرپذیری معرفی شده در (۴.۲) است.

تعریف ۲.۴. یکی از شاخص‌های ارزیابی نیکویی برازش مدل رگرسیون آماری، معیار میانگین خطای پیشگویی $^{\wedge}MPE$ است. تحت مفروضات مدل رگرسیونی (۱.۴) اگر $\hat{Y}_i$ مقدار پیشگویی

[^]Mean Predicted Error

Y_i باشد آنگاه MPE برابر است با:

$$MPE = \frac{1}{n} \sum_{i=1}^n d(Y_i, \hat{Y}_i) \quad (۸.۴)$$

که در این مقاله از d_2 (فاصله) وزن دار شده که در تعریف (۳.۳) آمده، استفاده کرده‌ایم.

۵. مثال کاربردی

در این بخش با بیان مثال عددی داده‌های سیمان پورتلند [۱۵] و انجام مقایسات بین چند روش، برتری عملکرد مدل ارائه شده در این مقاله روشن‌تر می‌شود. در این مثال اثر ترکیب سیمان پورتلند روی گرمای آزاد شده در حین سفت شدن آن، مورد بررسی قرار می‌گیرد.

جدول ۱: مجموعه داده‌های سیمان پورتلند و مقایسه چند روش مدلسازی رگرسیونی فازی بر اساس وزن‌ها و شاخص MPE

No.	x_1	x_2	x_3	(y_i, s_i)	W_1	W_2	W_3	d_1	d_2	d_3
1	7	26	6	(78.5, 6.9)	1	1	1	407.127	0.182	0.350
2	1	29	15	(74.3, 6.4)	1	0.884	0.007	260.154	4.056	14.082
3	11	56	8	(104.3, 9.4)	1	0.966	0.901	0.350	1.995	0.909
4	11	31	8	(87.6, 7.8)	1	0.961	0.985	67.456	2.246	0.377
5	7	52	6	(95.5, 8.6)	1	1	1	141.251	0.494	0.016
6	11	55	9	(109.2, 9.9)	1	0.718	0.0003	22.081	16.161	20.748
7	3	71	17	(102.7, 9.3)	1	1	1	38.456	1.120	0.000
8	1	31	22	(72.5, 6.2)	1	0.925	0.993	0.024	4.311	0.001
9	2	54	18	(93.3, 8.3)	1	0.853	0.006	5.463	5.510	14.399
10	21	47	4	(115.9, 10.6)	1	1	1	13.226	0.077	0.061
11	1	40	23	(68, 7.4)	1	0.0313	0	225.708	161.098	114.014
12	11	66	9	(113.3, 10.6)	1	0.99	0.8151	3.639	0.669	1.353
13	10	68	8	(109.4, 9.9)	1	0.861	0.12	7.944	6.990	5.512
MPE								91.76	15.762	13.2171

داده‌ها در جدول ۱ نمایش داده شده‌اند. این مسأله با ۱۳ مشاهده شامل ۳ متغیر مستقل شامل x_1 ، مقدار آلومینات تری کلسیم، x_2 ، مقدار سیلیکات تری کلسیم، x_3 ، مقدار آلومینو فریت تری کلسیم و متغیر وابسته (y_i, s_i) مقدار گرمای آزاد شده در هر گرم سیمان است. ۱۱-امین مشاهده متغیر وابسته را از $(83.8, 7.4)$ به $(68, 7.4)$ تغییر داده و آن را به‌عنوان یک داده دور افتاده در داده‌ها در نظر می‌گیریم. مجموعه داده‌ها در جدول ۱ نشان داده شده‌اند. در این جدول (W_1, d_1) ، (W_2, d_2) و (W_3, d_3) ، به ترتیب وزن‌ها و فاصله‌های بدست آمده از روش‌های کمترین قدم‌مطلق خطا، رحیمیان و همکاران [۱] و روش پیشنهادی است. جدول ۲ ضرایب مدل‌های رگرسیونی را با روش‌های مختلف اشاره شده در فوق، نشان می‌دهد. نتایج وزن‌های بدست آمده از شیوه‌های مختلف مدل‌های رگرسیونی، برتری روش پیشنهادی را در کم وزن دهی به داده‌های دور افتاده نشان می‌دهد. همانطور که مشاهده می‌کنید کمترین وزن را در بین ۴

جدول ۲: برآورد مدل رگرسیونی داده‌های مثال سیمان پورتلند به شیوه‌های مختلف

روش	برآورد رگرسیونی
کمترین قدرمطلق فازی	$\hat{y} = (0.28, 0.15) + (3.15, 0.34)x_1 + (0.96, 0.06)x_2 + (1.77, 0.17)x_3$
رحیمیان و همکاران	$\hat{y} = (49.13, 3.92) + (1.63, 0.15)x_1 + (0.66, 0.07)x_2 + (0.14, 0.01)x_3$
پیشنهادی	$\hat{y} = (47.91, 3.82) + (1.67, 0.16)x_1 + (0.67, 0.07)x_2 + (0.08, 0.001)x_3$

روش به داده یازدهم که از پرت بودن آن مطلع هستیم، روش پیشنهادی می دهد. همچنین روش پیشنهادی به داده هایی که تقریباً دور افتاده هستند نیز وزن کمتری را در بین مابقی روش ها می دهد. همچنین روش پیشنهادی کمترین مقدار MPE را در میان روش های دیگر داراست که تایید دیگری است بر مناسبت این روش استوار، زمانی که در بین داده ها، داده پرت وجود دارد.

۶. بحث و نتیجه گیری

هرگاه در بین مجموعه داده‌های فازی متغیر وابسته، داده دور افتاده وجود داشته باشد و متغیرها یا مشاهدات مربوط به آن‌ها دقیق نباشند، برآوردهای خوبی از روش‌های معمول FLS حاصل نخواهند شد و باید از روش‌های جایگزین رگرسیون استوار فازی استفاده کرد. در این مقاله یکی از این روش‌های جایگزین ارائه شد. در این روش که تعمیم روش رحیمیان و همکاران [۱] است، با توجه به معیارهای کم وزن دهی به داده‌های دور افتاده و هم معیار MPE ، روش تعمیم یافته پیشنهادی در این مقاله از مناسبت خوبی برخوردار بوده است که در یک مثال کاربردی به مقایسه این روش با روش‌های قبلی پرداخته شده است.

مراجع

۱. رحیمیان، ا.، ربیعی، م.ر. و شاهسونی، د. (۱۳۹۵). تحلیل رگرسیون استوار فازی با داده‌های خروجی و پارامترهای فازی، مجله اندیشه آماری، در دست چاپ.
2. Chang, P. T., and Lee, E. S. (1994). Ranking of fuzzy sets based on the concept of existence. Computers and Mathematics with Applications, 27(9-10), 1-21.
3. Chang, P. T., and Lee, E. S. (1996). A generalized fuzzy weighted least-squares regression. Fuzzy Sets and Systems, 82(3), 289-298.
4. Diamond, P. (1988). Fuzzy least squares. Information Sciences, 46(3), 141-157.
5. Draper, N. R., and Smith, H. (2014). Applied regression analysis. John Wiley and Sons.
6. Dubois, D. J. (1980). Fuzzy sets and systems: theory and applications (Vol. 144). Academic press.
7. Maronna, R.A., Martin, D.R., Yohai, V.J. (2006). Robust Statistics: Theory and Methods. John Wiley and Sons.
8. Hassanpour, H. and et al (2010). Fuzzy linear regression model with crisp coefficients: a goal programming approach. Iranian Journal of Fuzzy Systems, 7(2), 1-153.

9. Li, J. and etal. (2016). A new fuzzy regression model based on least absolute deviation. *Engineering Applications of Artificial Intelligence*, 52, 54-64.
10. Kim, K. and etal (2005). Least absolute deviation estimator in fuzzy regression. *J. Appl. Math. Comput*, 18, 649-656.
11. Taheri, S. M., and Kelkinnama, M. (2012). Fuzzy linear regression based on least absolutes deviations. *Iranian Journal of Fuzzy Systems*, 9(1), 121-140.
12. Tanaka, H. (1984). A formulation of fuzzy linear programming problem based on comparison of fuzzy numbers. *Control and cybernetics*, 3, 185-194.
13. Torabi, H., and Behboodian, J. (2007). Fuzzy least-absolutes estimates in linear models. *Communications in Statistics—Theory and Methods*, 36(10), 1935-1944.
14. Xu, R. (1991). A linear regression model in fuzzy environment. *Advance in Modelling Simulation*, 27(2).
15. Xu, R., Le, C. (2001). Multidimensional least-squares fitting with a fuzzy model. *Fuzzy Sets and Systems*, **119**(2), 215-223.



طراحی یک سیستم فازی-عصبی خود سازمانده با قابلیت یادگیری برخط برای شناسایی سیستم های دینامیک غیر خطی در حضور نویز

شیرین ریخته گرمشهد^۱، حمید طباطبائی^{۲*}

^۱گروه مهندسی کامپیوتر، واحد نیشابور، دانشگاه آزاد اسلامی، نیشابور، ایران

^{۲*}باشگاه پژوهشگران جوان و نخبگان، واحد مشهد، دانشگاه آزاد اسلامی، مشهد، ایران

چکیده. در این مقاله یک سیستم فازی-عصبی خود سازمانده با قابلیت یادگیری برخط برای شناسایی سیستم های دینامیک غیر خطی در حضور نویز ارائه شده است. این سیستم بدون هیچ نودی در لایه پنهان شبکه ی عصبی تابع شعاعی پایه شروع به کار می کند و در طی فرایند آموزش چنانچه معیارهای تولید قوانین برآورده شوند یک نرون قانون جدید به سیستم اضافه می شود. در فاز یادگیری پارامترها، الگوریتم آموزش حداقل مربعات بازگشتی با قابلیت یادگیری برخط برای افزایش سرعت همگرایی به کار گرفته شده است. بعد از فرایند ایجاد قانون جدید، کارایی سیستم، محاسبه شده و قوانینی که تأثیر کمتری در کارایی سیستم دارند هرس می شوند. نوآوری های اصلی در این مقاله عبارتند از ۱- استفاده از معیار جدید درجه ی تطبیق در فاز رشد قوانین؛ ۲- ارائه ی یک الگوریتم هرس جدید بر اساس چگالی که چگالی تعداد دفعاتی است که یک قانون آتش می شود، هر بار که یک الگو توسط یک قانون پوشش داده می شود به چگالی آن قانون یک واحد اضافه می شود در پایان قانونی که کمترین مقدار چگالی را داشته باشد از بین قوانین موجود هرس می شود؛ ۳- ایجاد قوانین فازی بدون استفاده از الگوریتم پس انتشار خطا؛ در پایان برای بررسی عملکرد، کارایی و دقت بالای سیستم یک مساله ی مبنای یک سیستم غیرخطی با درجه ی غیرخطی بالاست با مدل پیشنهادی شبیه سازی شده است. نتایج نشان می دهد که سیستم پیشنهادی حتی با وجود نویز نسبت به سایر روش ها می تواند با دقت بالاتر و ساختار فشرده تر سیستم های غیرخطی را مدل سازی کند.

مقدمه

یکی از مسائل کلیدی در شناسایی سیستم ها پیدا کردن یک ساختار مدل مناسب است که بتوان توسط آن مدل، رفتار فیزیکی سیستم را مشخص نمود. عموماً از این روش هایی برای مدل سازی و شناسایی سیستم ها مورد نیاز است. اغلب مسائلی که در دنیای واقعی وجود دارند به طور ذاتی غیر خطی هستند و با وجود اینکه سیستم های خطی به دست آوردن مدل ریاضی برای این سیستم ها کار دشواری است.

تحقیقات نشان می دهند که سیستم های فازی-عصبی در سال های اخیر به طور گسترده ای در شناسایی سیستم های غیر خطی مورد استفاده قرار گرفته اند. مزیت اصلی مدل سازی سیستم های فازی-عصبی توانایی تفسیر و قابلیت یادگیری این سیستم هاست. به طور کلی مدل سازی فازی-عصبی شامل دو فاز شناسایی ساختار و پارامترها است. در فاز شناسایی ساختار، قوانین فازی استخراج می شوند و در فاز شناسایی پارامترها پارامترهای قوانین فازی (مجموعه های

فازی و توابع عضویت فازی) به روز رسانی می شوند. چنانچه تعداد قوانین حاصل زیاد باشد، هزینه محاسباتی و زمان بیشتری در فرایند یادگیری سیستم مورد نیاز است و چنانچه تعداد قوانین کم باشد کارایی مطلوب به دست نخواهد آمد. از اینرو شناسایی ساختار نقش مهمی در مدل سازی سیستم های فازی-عصبی ایفا می کند. برای رویارویی با مشکلات فوق طراحی سیستم های فازی-عصبی خود سازماندهی مورد توجه قرار گرفت. در این سیستم ها نه تنها پارامترها بلکه ساختار نیز به طور تطبیقی در طی فرایند یادگیری تعیین می شوند و نیازی به دانش اولیه برای ایجاد قوانین فازی و مجموعه های فازی اولیه نیست.

اخیراً شبکه های فازی-عصبی خودسازماندهی گوناگونی برای شناسایی سیستم های غیر خطی به کار گرفته شده اند. Er و Wu در [۱] یک الگوریتم یادگیری خود سازماندهی سلسله مراتبی برای سیستم فازی-عصبی دینامیک (DFNN) ارائه کرده اند. DFNN قوانین فازی به فرم تاکاگی-سوگنو را ایجاد می کند که در این قوانین عرض توابع عضویت گوسی برای تمام متغیرهای ورودی یکسان است که این امر واقع بینانه نیست مخصوصاً زمانی که متغیرهای ورودی بازه های مختلف عملکردی داشته باشند. از این رو نسخه ای اصلاح شده ای این روش در [۲] به نام سیستم GDFNN ارائه شد که در این سیستم برای هر تابع عضویت گوسی داخل هر نرون تابع شعاعی پایه، بردار عرض متناسب معرفی شده است. اگر چه در این روش نیز الگوریتم یادگیری سلسله مراتبی برای آموزش GDFNN وابسته به کل داده های آموزش می باشد از این روش GDFNN برای یادگیری برون خطی مناسب است.

Wang در [۳] یک شبکه ی فازی-عصبی خودسازماندهی برخط سریع با دقت بالا (FAOS-PFNN) ارائه کرده است که این سیستم ابتدا بدون هیچ قانونی شروع به کار می کند و بر اساس یک معیار رشد بر حسب نیاز قوانین به لایه ی پنهان در طی فرایند یادگیری اضافه می شوند. متأسفانه در FAOS-PFNN امکان هرس نرون های اضافی (قوانین) لایه ی پنهان وجود ندارد از سوی دیگر تحلیلی برای همگرایی FAOS-PFNN ارائه نشده است. نسخه ی تعمیم یافته ی FAOS-PFNN مدل GEBF-OSFNN [۴] می باشد که در این مدل از توابع عضویت گوسی نامتقارن به دلیل افزایش انعطاف پذیری و قابلیت تفسیر بالاتر فضای ورودی استفاده شده است. به منظور ایجاد قوانین فازی توسط یک معیار رشد از یک مجموعه داده ی بزرگ، یک شبکه ی فازی خودسازماندهی (SOFLMS) توسط Rubio در [۵] ارائه شده است که توسط الگوریتم آموزش حداقل مربعات اصلاح شده پارامترهای سیستم فازی به روزرسانی می شوند. در SOFLMS از الگوریتم نزدیک ترین همسایه برای ایجاد قوانین استفاده شده است. به این معنی که یک قانون جدید در صورتی ایجاد می شود که فاصله ی بین الگوی ورودی از مراکز نرون های تابع شعاعی پایه موجود از یک حد از پیش تعیین شده ای بیشتر باشد. به طور هم زمان از یک الگوریتم هرس برای هرس قوانین اضافی استفاده شده است. اگرچه قوانین فازی در SOFLMS فقط بر اساس الگوی جدید تعیین می شوند اما ایراد این روش این است که در مرحله ی شناسایی ساختار تنها یک نرون می تواند حذف شود.

در بعضی از روش ها از الگوریتم های تکاملی مانند الگوریتم ژنتیک [۶] و [۷]، برای به دست آوردن تعداد قوانین بهینه استفاده شده است. در روش های مبتنی بر الگوریتم ژنتیک تعداد قوانین فازی و همچنین پارامترهای توابع عضویت در غالب یک کروموزوم پیچیده کدگذاری می شوند. در پایان قانونی که قسمت مقدم یا قسمت تالی صفر داشته باشد در مجموعه قوانین فازی نهایی وجود نخواهد داشت. هر چند که تعداد مجموعه های فازی برای هر متغیر ورودی مبتنی بر دانش از پیش تعیین شده ای است اما با این روش ماکزیمم تعداد قوانین به دست می آید.

RSMode [۸] یک الگوریتم یادگیری برای یک سیستم فازی-عصبی خودسازماندهی است که در این الگوریتم از روش چند زیر جمعیتی^۱ در فاز یادگیری قوانین استفاده شده است که برای هر زیر جمعیت به طور جداگانه یک کروموزوم در

¹ Multi-subpopulation

نظر گرفته شده است که هر کروموزوم بیانگر یک قانون است و به طور جداگانه مورد ارزیابی قرار می گیرند. پارامترها توسط الگوریتم تکاملی تفاضلی بهبود داده شده تعیین می شوند.

SOME [۹] یک الگوریتم یادگیری برای شبکه ی فازی-عصبی تاکاگی سوگنو می باشد که در این الگوریتم از گروه مبتنی بر تکامل هم زیستی^۱ استفاده شده است که در هر گروه فقط یک قانون فازی نمایش داده می شود. در فاز یادگیری ساختار از یک الگوریتم خود سازماندهی دو مرحله ای و در فاز یادگیری پارامترها از یک روش انتخاب مبتنی بر داده کاوی^۲ استفاده شده است. بزرگترین نقطه ضعف الگوریتم های مبتنی بر الگوریتم ژنتیک زمان آموزش طولانی است.

روش هایی در [۱۰] و [۱۱] برای ایجاد قوانین فازی مبتنی بر ماشین بردار پشتیبان ارائه شده است. در این مقالات از ماشین بردار پشتیبان برای ایجاد خودکار قوانین فازی استفاده شده است که تعداد قوانین برابر با تعداد ماشین بردار پشتیبان هاست. مهم ترین مسأله در یادگیری ماشین بردار پشتیبان، مسأله ی نرون های مرده^۳ می باشد. این که چگونه می توان توسط ماشین بردار پشتیبان تعداد بهینه ای قوانین فازی را برای سیستم های فازی-عصبی^۴ تعیین نمود هنوز یک مسأله ی چالش برانگیز است. در [۱۲] نیز از الگوریتم آلفوریتم ژنتیک برای شناسایی قوانین فازی استفاده شده است و تنظیم وزن های قوانین به صورت جداگانه و در مراحل مختلف صورت می گیرد که این امر منجر به کاهش طول کروموزوم ها و کوچک شدن فضای جستجو می گردد. بزرگترین نقطه ضعف الگوریتم های مبتنی بر آلفوریتم ژنتیک زمان آموزش طولانی است. در [۱۳] از یک الگوریتم خوشه بندی جدید برای بدست آوردن تعداد قوانین فازی بهینه و یک الگوریتم هرس مبتنی بر روش hebbian برای کاهش تعداد قوانین فازی استفاده شده است. در [۱۴] از مقادیر فازی نوع دو برای مقداره های وزن ها استفاده شده است و الگوریتم کارنیک مندل بهبود یافته شده برای تنظیم پارامترها و یک الگوریتم خوشه بندی فازی نوع دو برای پیدا کردن نرون ها در لایه پنهان به کار گرفته شده است

در این نوشتار یک سیستم فازی-عصبی خود سازمانده با قابلیت یادگیری برخط برای شناسایی سیستم های دینامیک غیر خطی در حضور نویز ارائه شده است. ابتدا در این سیستم هیچ نودی در لایه ی پنهان وجود ندارد و در طی فرایند آموزش چنانچه معیارهای تولید قوانین برآورده شوند نرون تابع شعاعی پایه^۵ به لایه ی پنهان اضافه می شود. پارامترهای قانون جدید به فرم تاکاگی-سوگنو با استفاده از الگوریتم آموزش حداقل مربعات بازگشتی^۶ برای قابلیت یادگیری برخط و افزایش سرعت همگرایی، تخمین زده می شوند. بعد از فرایند ایجاد قانون جدید، کارایی سیستم، محاسبه شده و قوانینی که تأثیر کمتری در کارایی سیستم دارند هرس می شوند.

نوآوری های اصلی این مقاله عبارت است از:

- (۱) استفاده از دو معیار جدید در فاز رشد قوانین: درجه ی تطبیق و معیار متداول فاصله؛
- (۲) ارائه ی یک الگوریتم هرس جدید بر اساس چگالی: چگالی تعداد دفعاتی است که یک قانون آتش می شود، هر بار که یک الگو توسط یک قانون پوشش داده می شود به چگالی آن قانون یک واحد اضافه می شود در پایان قانونی که کمترین مقدار چگالی را داشته باشد از بین قوانین موجود هرس می شود؛
- (۳) ترکیب توابع عضویت مشابه و تغییر عرض آنها به منظور افزایش کارایی سیستم؛
- (۴) ایجاد قوانین فازی بدون استفاده از الگوریتم پس انتشار خطا: بر اساس الگوریتم آموزش حداقل مربعات بازگشتی؛

^۱ group-based symbiotic evolution

^۲ Selection strategy Based on Data mining

^۳ Support Vector Machine

^۴ dead neurons problem

^۵ Neuro Fuzzy Systems

^۶ Radial Basis Function

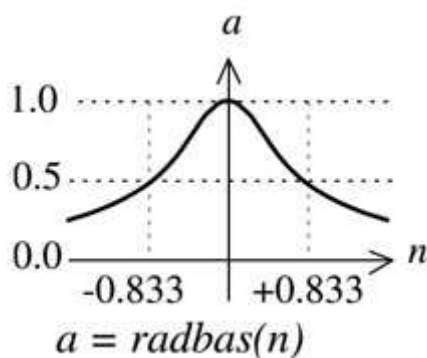
^۷ Recursive Least Square

نتایج شبیه سازی با بررسی عملکرد سیستم در حضور نویز روی دو مسأله ی مینا، شامل شناسایی سیستم های دینامیک غیرخطی حاکی از دقت بالاتر و ساختار فشرده تر الگوریتم پیشنهادی نسبت به سایر روش هاست. این مقاله در ۵ بخش سازماندهی شده است. در ادامه در بخش ۲ بعد از مقدمه، مفاهیم مقدماتی ارائه شده است. ساختار OSO-NFS که معادل با قوانین فازی نوع تاکاگی- سوگنو است و الگوریتم یادگیری شامل الگوریتم های رشد و هرس شبکه به تفصیل در بخش ۳ شرح داده شده است، در بخش ۴ شبیه سازی الگوریتم پیشنهادی ارائه شده است و در نهایت ارزیابی نتایج و مقایسه ی نتایج با سایر روش ها در بخش ۵ صورت گرفته است.

مفاهیم مقدماتی

شبکه های تابع شعاعی پایه:

در حوزه ی مدلسازی ریاضی، RBF یک شبکه عصبی مصنوعی هستش که از توابع پایه ای شعاعی به عنوان توابع فعالیت استفاده می کند [3] خروجی این شبکه یک ترکیب خطی از توابع پایه ی شعاعی برای پارامترهای ورودی و نرونهاست. این شبکه ها در تابع تقریب، پیشبینی سری های زمانی، کلاس بندی و کنترل سیستم مورد استفاده قرار می گیرند. شبکه های شعاع مبنا به نسبت شبکه پس انتشار نیاز به نرون های بیشتری دارند، اما حسن آن ها در زمان طراحی کوتاه تر آن ها نسبت به شبکه های استاندارد پس انتشار می باشد. این شبکه ها زمانی که بردارهای آموزشی بسیار زیاد می باشند، کارایی قابل ملاحظه ای از خود نشان می دهند.



شکل ۲- تابع انتقال شعاعی

این تابع با ورودی صفر دارای حداکثر مقدار خود یعنی ۱ می باشد. به این ترتیب با کاهش فاصله بین p و W ، خروجی تابع انتقال شعاعی افزایش می یابد. بنابراین نرون شعاع مبنا به عنوان یک تشخیص دهنده یکسان شدن W و p عمل می کند. به این معنی که زمانی که p و W یکسان شوند، تابع انتقال شعاعی، خروجی ۱ را تولید می کند. مقدار بایاس b ، تنظیم میزان حساسیت نرون شعاع مبنا را میسر می سازد. به عنوان مثال اگر یک نرون دارای مقدار بایاس 0.1 باشد، خروجی نرون به ازای فاصله 0.8326 بین p و W برابر 0.5 خواهد بود. زیرا با ضرب بایاس در فاصله $0.1 \times 0.8326 = 0.08326$ ، مقدار 0.8326 حاصل می شود که خروجی تابع شعاع مبنا برای آن 0.5 می باشد. حال اگر بایاس 0.2 به ازای فاصله 4.163 باشد، خروجی 0.5 حاصل می شود و به این معنی است که حساسیت نرون بیشتر شده است.

روش پیشنهادی OSO-NFS

در این بخش ساختار و معماری ۵ لایه ی OSO-NFS که معادل با قوانین فازی نوع تاکاگی- سوگنو است و هم چنین الگوریتم یادگیری شامل الگوریتم های رشد و هرس شبکه و ترکیب توابع عضویت مشابه به منظور افزایش کارایی و ساده سازی و کاهش پیچیدگی سیستم ارائه شده است.

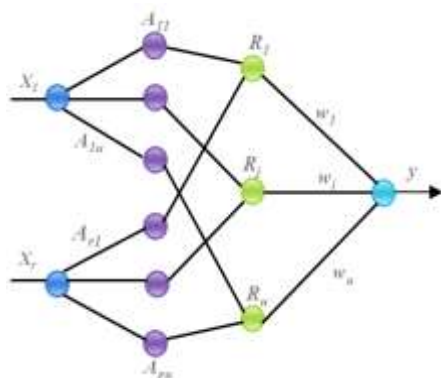
معماری OSO-NFS

در این بخش معماری OSO-NFS در شکل ۳ بیان شده است. یک شبکه ی چهار لایه ی مشابه با [۱] و [۲] برای پیاده سازی قوانین فازی نوع تاکاگی - سوگنو به فرم زیر ارائه شده است: فرض می کنیم تعداد متغیر های ورودی r است.

Rulej: IF x_1 is A_{1j} and ... and x_r is A_{rj} THEN y is W_j $j = 1, 2, \dots, u$ (2)

$$W_j = \alpha_{0j} + \alpha_{1j}x_1 + \dots + \alpha_{rj}x_r$$

در رابطه ی فوق $A_{1j}, \dots, A_{rj}$ مجموعه های فازی قسمت مقدم قوانین (توابع عضویت گوسی) $\alpha_{0j}, \alpha_{1j}, \dots, \alpha_{rj}$ پارامترهای قسمت تالی و وزن های شبکه هستند



شکل 3: معماری OSO-NFS [۱]

لایه ی اول: هر نود در این لایه نشان دهنده ی یک متغیر ورودی است. $x_i (i=1, 2, \dots, r)$
 لایه ی دوم: هر نود در این لایه معرف تابع عضویت است. هر متغیر ورودی در لایه ی اول U مجموعه ی فازی دارد A_{ij} ($j = 1, 2, \dots, u$)

$$\mu_{ij}(x_i) = \exp\left(-\frac{(x_i - c_{ij})^2}{\sigma_{ij}^2}\right), i = 1, 2, \dots, r, \\ j = 1, 2, \dots, u \quad (3)$$

که μ_{ij} ، زامین تابع عضویت متغیر x_i و c_{ij} و σ_{ij} به ترتیب مرکز و عرض این توابع عضویت هستند.
 لایه ی سوم: هر نود در این لایه نود قانون نامیده می شود که قسمت اگر قوانین فازی را نشان می دهد (ت-نرم قسمت مقدم قوانین) خروجی z امین قانون ($j = 1, 2, \dots, u$) R_j به فرم زیر است:

$$\varphi_j(x_1, x_2, \dots, x_r) = \exp\left(-\sum_{i=1}^r \frac{(x_i - c_{ij})^2}{\sigma_{ij}^2}\right) \\ j = 1, 2, \dots, u \quad (4)$$

لایه ی چهارم: هر نود در این لایه یک نود خروجی است و خروجی این نود از حاصل جمع سیگنال های وارد شده به آن به دست می آید.

$$y(x_1, x_2, \dots, x_r) = \sum_{j=1}^u w_j \phi_j \quad (5)$$

در شکل 3 لایه ی خروجی دارای یک نود است برای نمایش ساده ترکه معرف سیستم چند ورودی تک خروجی است. قابل ذکر است که نتایج می توانند به یک سیستم چند ورودی چند خروجی بسط داده شوند.

الگوریتم یادگیری در OSO-NFS

دو فاز مهم در طراحی سیستم های فازی-عصبی عبارتند از فاز یادگیری پارامترها (تنظیم پارامترهای قسمت مقدم و تالی قوانین و وزن های شبکه) و یادگیری ساختار که همان فاز تولید قوانین بهینه است. در این بخش الگوریتم هایی به این منظور معرفی شده است.

معیار های تولید نرون (قوانین) در لایه ی پنهان

در این بخش دو معیار جدید عبارت از معیار فاصله و معیار درجه ی تطبیق برای ایجاد قوانین جدید ارائه شده است. این سیستم در ابتدا بدون هیچ نودی در لایه پنهان شروع به کار می کند و در طی فرایند یادگیری چنانچه معیار های تولید قوانین برآورده شود قوانین به لایه ی پنهان اضافه می شوند.

معیار فاصله:

تعریف ۱: [6] ε -Completeness

برای هر متغیر ورودی حداقل یک قانون فازی وجود دارد به نحوی که درجه ی تطبیق (میزان آتش شدن) کمتر از ε نیست.

نکته: در کاربردهای فازی معمولاً حداقل مقدار $\varepsilon = 0.5$ در نظر گرفته می شود. اگر یک الگوی جدید شرط ε -Completeness را برآورده کند OSO-NFS قانون جدید ایجاد نخواهد کرد و تنها پارامترهای الگوی جدید به روز رسانی می شوند. به این ترتیب که با مشاهده ی زوج ورودی $(X^k, t^k), k=1,2,\dots,n$ - M فاصله $(md^k(j))$ بین X^k و مراکز واحدهای RBF موجود $(c_j(j=1,2,\dots,u))$ طبق رابطه ی زیر محاسبه می شود:

$$md(j) = \sqrt{(X - C_j)^T \sum_j -1 (X - C_j)} \quad (6)$$

که

$$C_j = (c_{1j}, c_{2j}, \dots, c_{rj})^T \text{ و } X = (x_1 \dots x_r)^T$$

$$\sum_j -1 = \begin{bmatrix} \frac{1}{\sigma_{1j}^2} & 0 & \dots & 0 \\ 0 & \frac{1}{\sigma_{2j}^2} & 0 & 0 \\ 0 & 0 & \ddots & 0 \\ 0 & \dots & 0 & \frac{1}{\sigma_{rj}^2} \end{bmatrix} \quad j = 1, 2, \dots, u \quad (7)$$

سپس پیدا می کنیم:

$$J = \arg \min_{1 \leq j \leq u} (md^k(j)) \quad (8)$$

$$md_{\min}^k = md^k(J) > k_d \quad (9)$$

که در رابطه ی فوق k_d یک مقدار از پیش تعیین شده ای است که در طی فرایند یادگیری کاهش می یابد. رابطه ی (7) بر این امر دلالت دارد که چنانچه سیستم موجود توسط ϵ -Completeness برآورده نشود یک قانون جدید باید تولید شود. k_d طبق رابطه ی زیر به دست می آید:

$$\begin{cases} d_{\max} = \sqrt{\ln(1/\epsilon_{\min})} & 1 < k < n/3 \\ \max[d_{\max} \times \gamma^k, d_{\min}] & n/3 \leq k \leq 2n/3 \\ d_{\min} = \sqrt{\ln(1/\epsilon_{\max})} & 2n/3 < k \leq n \end{cases} \quad (10)$$

که $\gamma \in (0,1)$ ثابت کاهش نامیده می شود و طبق رابطه ی زیر به دست می آید :

$$\gamma = \left(\frac{d_{\min}}{d_{\max}}\right)^{\frac{3}{n}} \quad (11)$$

معیار درجه تطبیق:

با ورود الگوی i ام ($i=1,2,\dots,n$) درجه ی تطبیق این داده ی جدید با هر نرون قانون در لایه ی سوم طبق رابطه ی (4) محاسبه خواهد شد. چنانچه ماکزیمم این مقدار از یک حد آستانه ای کمتر باشد به این معنا است که این داده نمی تواند توسط قوانین موجود پوشش داده شود و در این صورت یک قانون جدید باید ایجاد شود.

$$\Phi_{\max} = \max(\Phi_j(x)), j = 1, \dots, u \quad (12)$$

که u تعداد نرون ها در لایه ی پنهان است. چنانچه $\Phi_{\max} < \mu_r$ یک قانون جدید باید ایجاد شود.

الگوریتم هرس شبکه (حذف قوانین اضافی)

بعضی اوقات ممکن است یک نرون در ابتدا فعال شود اما در پایان سهم کوچکی در تولید خروجی سیستم داشته باشد. علاوه بر این هرس شبکه برای سیستم های غیر خطی دینامیک متغیر با زمان امری ضروری است. اگر نرون های غیر فعال در لایه ی پنهان قابل شناسایی و حذف باشند در طی فرایند یادگیری می توان به شبکه ای با توپولوژی فشرده تر دست پیدا کرد که این امر در نهایت منجر به کاهش زمان آموزش خواهد شد. در این بخش یک الگوریتم جدید بر اساس چگالی ارائه خواهد شد. چگالی تعداد دفعاتی است که هر قانون آتش شده است. زمانی که یک قانون جدید ایجاد می شود چگالی آن یک قرار داده می شود. $d_{M+1}=1$ که M تعداد قوانین موجود است. با ورود هر الگوی جدید چنانچه این داده توسط یکی از قوانین موجود در سیستم پوشش داده شود چگالی آن قانون یک واحد اضافه خواهد شد. در پایان قانون کم اهمیت قانونی است که کمترین مقدار چگالی را داشته باشد. بعد از چند دوره $\Delta(L)$ چنانچه مقدار چگالی یک قانون از حد آستانه ای کمتر باشد آن قانون حذف خواهد شد.

$$d_{\min} = \min(d_j), 1 \leq j \leq u \quad (13)$$

اگر $u > 2$ و $d_{\min} \leq d_u$ این قانون حذف خواهد شد که d_u مینیمم چگالی انتخاب شده به نام پارامتر آستانه است.

تنظیم پارامترهای قسمت مقدم

بعد از رشد یک نرون گام بعد تعیین پارامترهای قسمت تالی است پارامترهای نود RBF طبق روابط زیر تعیین می شوند:

$$C_i = X_i \quad (13)$$

$$\sigma_i = \frac{\max\{|c_i - c_{i-1}|, |c_i - c_{i+1}|\}}{\sqrt{\ln(\frac{1}{\epsilon})}} \quad (14)$$

در رابطه ی فوق c_{i-1} و c_{i+1} مراکز دو تابع عضویت همسایه هستند.

تنظیم پارامترهای قسمت تالی

فرض کنید u قانون فازی برای n نمونه الگوی ورودی خروجی برای r متغیر ورودی تولید شده باشد با نوشتن رابطه ی (5) به فرم ماتریس داریم:

$$W\varphi = Y \quad (15)$$

در رابطه ی فوق:

$$Y \in R^n \quad \text{و} \quad \varphi \in R^{u(r+1) \times n}, \quad W \in R^{u(r+1)}$$

اگر خروجی مطلوب سیستم $T = (t_1, t_2, \dots, t_n) \in R^n$ باشد هدف تعیین پارامترهای بهینه ی W^* به نحوی است که رابطه ی خطی $\|W\varphi - T\|_2$ حداقل شود با استفاده از روش حداقل مربعات خطی مقدار بهینه ی W^* از رابطه ی زیر به دست می آید:

$$W^*(k) = T(k)(\phi(k)^T \phi(k))^{-1} \phi(k)^T. \quad (16)$$

که در رابطه ی فوق

$$T(k) = [d^T(1) \ d^T(2) \ \dots \ d^T(n)]^T$$

$$\phi(k) = [\varphi(1) \ \varphi(2) \ \dots \ \varphi(n)]^T$$

برای ساده سازی داریم

$$P(k) = [(\phi(k)^T \phi(k))]. \quad (17)$$

با استفاده از ماتریس P معادلات نهایی در الگوریتم حداقل مربعات بازگشتی به صورت زیر می باشد

$$L(k) = P(k)\varphi(k) = P(k-1)\varphi(k)[1 + \varphi^T(k)P(k-1)\varphi(k)]^{-1}$$

$$P(k) = [I - L(k)\varphi^T(k)]P(k-1)$$

$$W^*(k) = W^*(k-1) + P(k)\phi(k)[T(k) - \varphi^T(k)W^*(k-1)].$$

(18)

نتایج نشان می دهد چون ساختار OSO-NFS بعد از فرایند یادگیری ساختار بسیار فشرده است ابعاد W و φ کوچک است.

و به کمک روش RLS با انجام محاسبات ساده می توان پارامترهای قسمت تالی را به نحو مناسبی تخمین زد.

۴. شبیه سازی و ارزیابی نتایج

در این بخش فرایند شناسایی یک سیستم غیر خطی با درجه ی غیرخطی بودن بالا به منظور نمایش کارایی OSO-NFS مورد بررسی قرار گرفته است. برای نمایش عملکرد مناسب سیستم پیشنهادی فرایند شناسایی در دو حالت با حضور نویز و بدون نویز بررسی شده است. و نتایج با الگوریتم های متعددی مقایسه شده است. در این مقاله برای پیاده سازی سیستم پیشنهادی از نرم افزار MATLAB R2011a استفاده شده است. شبیه سازی ها بر روی CPU پنتیوم ۳/۲ گیگا هرتز با ۴ گیگا بایت حجم اجرا شده است.

مثال- شناسایی یک سیستم غیر خطی ورودی محدود-خروجی محدود

سیستم زیر یکی از مسائل مبناست که در [۱۳، ۱۴، ۱۵] به کار گرفته شده است که با رابطه زیر نشان داده می شود:

$$y(k) = u(k)^3 + \frac{y(k-1)}{1 + y(k-1)^2} \quad (19)$$

که $u(k)$ ورودی فعلی و $y(k)$ و $y(k-1)$ به ترتیب خروجی لحظه ی فعلی و قبلی سیستم هستند. $u(k)$ از رابطه ی زیر محاسبه می شود:

$$u(k) = \begin{cases} -0.7 + \frac{\text{mod}(k, 5)}{40}, & k \leq 80 \\ \text{rands}(1, 1), & 80 < k \leq 130 \\ 0.7 - \frac{\text{mod}(k, 180)}{180}, & 130 < k \leq 250 \\ 0.6 \cos\left(\frac{\pi k}{50}\right), & 250 < k \leq 400 \end{cases} \quad (20)$$

تابع $\text{mod}(x, y)$ باقیمانده x/y است. در این مثال ورودی $u(k)$ به ازای $k=1 \dots 400$ به سیستم اعمال می شود.

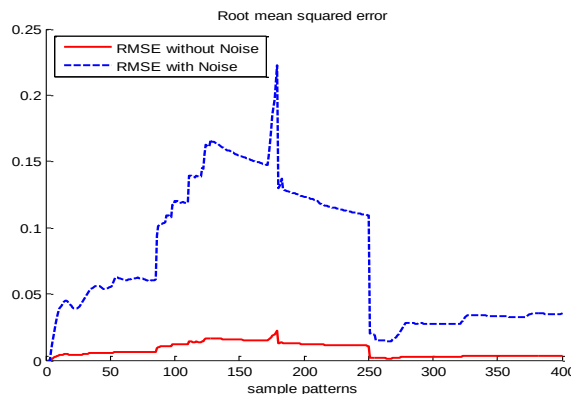
در فاز تست 40° داده در بازه ی $k=1$ تا 400 به صورت زوج ورودی خروجی در نظر گرفته شده است.

بهینه ترین مقادیر اولیه ی پارامترها در طی فرایند یادگیری با استفاده از روش تحلیل پارامترها و سعی و خطا به صورت زیر در نظر گرفته می شود:

$$\begin{aligned} \varepsilon_{\min} &= 0.2, \varepsilon_{\max} = 2, e_{\min} = 0.02 \\ e_{\max} &= 0.9, k_{mf} = 0.6, k_s = 0.6, k_{err} = 0.007, d_u = 4 \end{aligned}$$

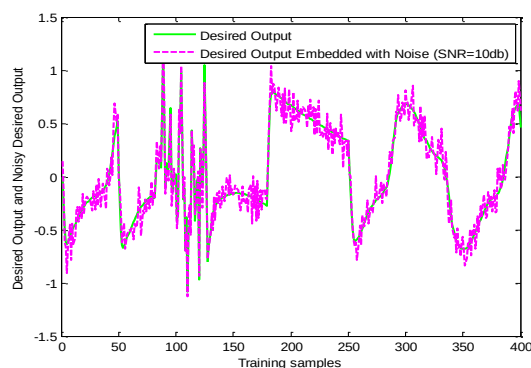
نتایج شبیه سازی در طی فرایند یادگیری برخط در شکل های ۴-۱۰ شده است. نتایج نشان می دهد که OSO-NFS می تواند زمانی که داده های آموزش به صورت برخط به سیستم وارد می شوند (الگو به الگو) سیستم غیر خطی را در حالی که خطای خروجی و خطای جذر میانگین مربعات^۱ به سمت صفر میل می کنند، شناسایی کند.

نتایج شبیه سازی مثال در شکل های ۴-۱۰ نشان داده شده است. مقایسه ی نتایج بین OSO-NFS و سایر الگوریتم های معروف در جداول ۱ و ۲ آورده شده است. نتایج نشان می دهد که سیستم پیشنهادی در شناسایی سیستم های غیر خطی با درجه غیرخطی بودن بالا نسبت به شبکه های عصبی و شبکه های فازی-عصبی خود سازمانده نتایج قابل قبول تری ارائه می دهند. سیستم پیشنهادی با ساختار فشرده تر و تعداد قوانین کمتر توانسته است با خطای کمتری سیستم غیر خطی را شناسایی کند.



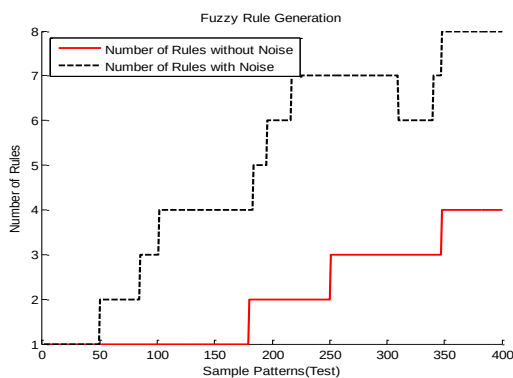
شکل ۴: ریشه‌ی میانگین مربعات خطا (RMSE) در حضور نویز و بدون نویز

همان طور که در شکل ۴ واضح است ریشه‌ی میانگین مربعات خطا در طی فرایند یادگیری برخط در هر دو حالت نویزی و بدون نویز در حال کاهش است



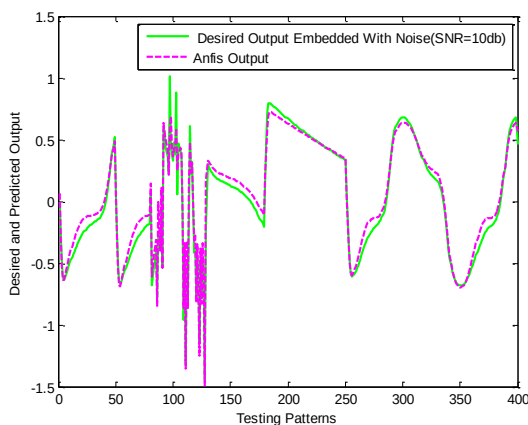
شکل ۵: نمایش سیگنال اصلی و سیگنال اصلی+ نویز گوسی سفید

در شکل ۵ سیگنال اصلی بدون نویز و همراه با نویز گوسی سفید با توان نویز ۱۰ db نشان داده شده است.



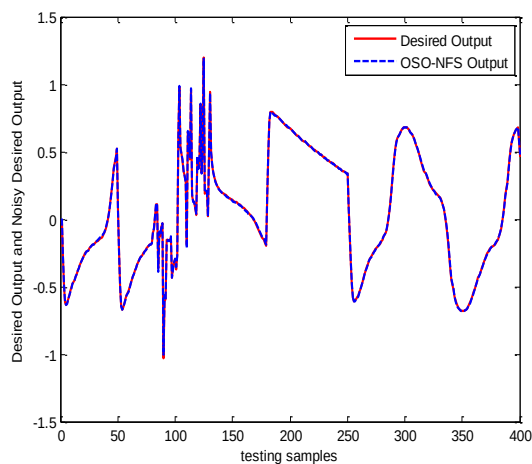
شکل ۶: رشد نرون‌های قانون‌نوس-OSO در حضور نویز و بدون نویز

رشد نرون‌های قانون در شکل ۶ در دو حالت نویزی و بدون نویز نمایش داده شده است. سیستم پیشنهادی با ۴ قانون در حالت بدون نویز و ۸ قانون در حالت نویزی توانسته سیستم را با کمترین میزان خطا مدلسازی کند.



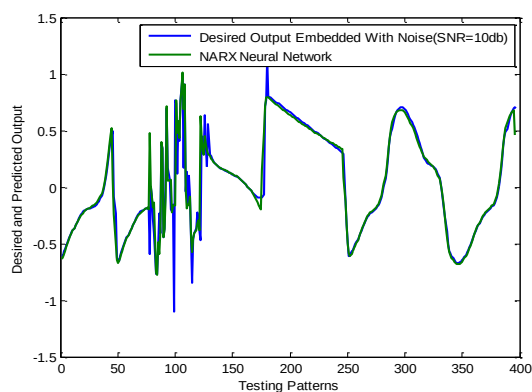
شکل ۷: مقایسه‌ی خروجی واقعی و خروجی ANFIS در حضور نویز $\text{SNR}=10\text{db}$ و $\text{noise power}=0$

شکل ۷ مقدار خروجی واقعی و خروجی ANFIS را در مدل‌سازی سیستم مبنا نشان داده است.



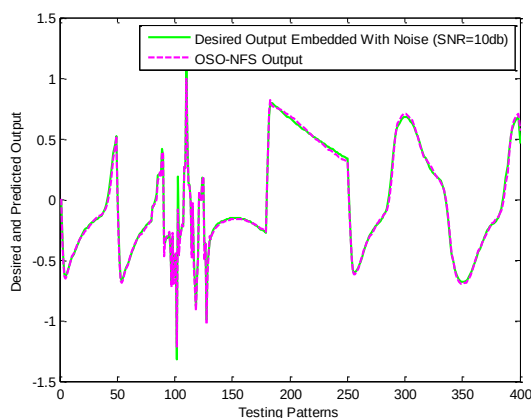
شکل ۸: مقایسه‌ی خروجی واقعی و خروجی OSO-NFS بدون نویز

شکل ۸ مقدار خروجی واقعی و خروجی سیستم پیشنهادی را در مدل‌سازی سیستم مبنا بدون نویز نشان داده است.



شکل ۹: مقایسه‌ی خروجی واقعی و خروجی شبکه عصبی NARX در حضور نویز $\text{SNR}=10\text{db}$ و $\text{noise power}=0$

شکل ۹ مقدار خروجی واقعی و خروجی سیستم شبکه ی عصبی NARX را در مدل سازی سیستم مبنا نشان داده است.



شکل ۱۰: مقایسه ی خروجی واقعی و خروجی OSO-NFS در حضور نویز SNR=10db و noise power=0
شکل ۱۰ مقدار خروجی واقعی و خروجی سیستم پیشنهادی را در مدل سازی سیستم مبنا همراه با نویز نشان داده است.

جدول ۱- مقایسه ی نتایج OSO-NFS مثال ۲ با سایر روش ها بدون نویز

Algorithm	سیستم غیرخطی ورودی محدود-خروجی محدود	
	RMSE	No. of final RBF unit
T2RBFN[19]	0.0032	2
NARX Neural Network	0.0396	10
ANFIS	0.0207	5
OSO-NFS(without noise)	0.0024	4

نتایج جدول ۱ نشان می دهد شبکه های فازی-عصبی در مقایسه با شبکه های عصبی چند لایه NARX قدرت بیشتری در شناسایی سیستم های غیر خطی دارند و همان طور که آشکار است نسبت به ANFIS که یک شبکه فازی عصبی با ساختار ثابت است شبکه های فازی عصبی خودسازمانده با ساختار فشرده تر و میزان خطای کمتر توانایی مدل سازی سیستم های غیرخطی را دارا می باشند.

جدول ۲-مقایسه ی نتایج OSO-NFS مثال ۲ روش ANFIS و NARX در حضور نویز با مقادیر مختلف نرخ سیگنال به نویز توان نویز

Algorithm	سیستم غیرخطی ورودی محدود-خروجی محدود		
	Rule number	SNR(db), noise power(db)	RMSE
OSO-NFS	9	25,1	0.1812

Algorithm	سیستم غیرخطی ورودی محدود-خروجی محدود		
	Rule number	SNR(db), noise power(db)	RMSE
	10	25,2	0.1998
	8	20,1	0.1925
	9	20,2	0.2001
NARX	16	25,1	0.2545
	16	25,2	0.2671
	16	20,1	0.2989
	16	20,2	0.3045
ANFIS	12	25,1	0.2178
	11	25,2	0.2334
	13	20,1	0.2457
	14	20,2	0.2532

در جدول ۲ با اعمال نویز با توان نویزها و نرخ سیگنال به نویزهای متفاوت، با استفاده از هر سه روش سیستم مبنا مدلسازی شده است که بررسی این نتایج نشان می دهد که سیستم پیشنهادی عملکرد بهتری دارد و با ساختار فشرده تر و میزان خطای کمتری توانسته سیستم مبنا را مدلسازی نماید.

۵- نتیجه گیری

در این مقاله یک سیستم فازی-عصبی خود سازمانده با قابلیت یادگیری برخط برای شناسایی سیستم های دینامیک غیر خطی در حضور نویز ارائه شده است. ابتدا در این سیستم هیچ نودی در لایه ی پنهان وجود ندارد و در طی فرایند آموزش چنانچه معیارهای تولید قوانین برآورده شوند نرون تابع شعاعی پایه به لایه ی پنهان اضافه می شود. پارامترهای قانون جدید به فرم تاکاگی-سوگنو با استفاده از الگوریتم آموزش حداقل مربعات بازگشتی برای قابلیت یادگیری برخط و افزایش سرعت همگرایی، تخمین زده می شوند. بعد از فرایند ایجاد قانون جدید، کارایی سیستم، محاسبه شده و قوانینی که تأثیر کمتری در کارایی سیستم دارند هرس می شوند. ویژگی های اصلی در این مقاله عبارتند از استفاده از معیار جدید درجه ی تطبیق و معیار متداول فاصله در فاز رشد قوانین؛ ارائه ی یک الگوریتم هرس جدید بر اساس چگالی که چگالی تعداد دفعاتی است که یک قانون آتش می شود، هر بار که یک الگو توسط یک قانون پوشش داده می شود به چگالی آن قانون یک واحد اضافه می شود در پایان قانونی که کمترین مقدار چگالی را داشته باشد از بین قوانین موجود هرس می شود؛ ایجاد قوانین فازی بدون استفاده از الگوریتم پس انتشار خطا؛ همان طور که نتایج نشان داد سیستم پیشنهادی با دقت بالاتر و ساختار فشرده تر نسبت به سایر الگوریتم ها توانایی مدلسازی سیستم های غیر خطی را دارد. این در شرایطی است که خروجی سیستم غیر خطی بدون نویز و یا حتی همراه با نویز باشد. با بررسی عملکرد سیستم در حضور نویز روی یک مسأله ی مبنا، شامل شناسایی یک سیستم دینامیک غیرخطی برتری روش پیشنهادی نسبت به روش های دیگر نشان داده شد.

مراجع:

- [1] S. Wu and M. J. Er, "Dynamic fuzzy neural networks—A novel approach to function approximation," IEEE Trans. Syst., Man, Cybern. B, Cybern., vol. 30, no. 2, pp. 358–364, Apr. 2000.
- [2] S. Wu, M. J. Er, and Y. Ho, "A fast approach for automatic generation of fuzzy rules by generalized dynamic fuzzy neural networks," IEEE Trans. Fuzzy Syst., vol. 9, no. 4, pp. 578–594, Aug. 2001.
- [3] N. Wang, M. J. Er, and X. Y. Meng, "A fast and accurate online self-organizing scheme for parsimonious fuzzy neural networks," Neurocomput., vol. 72, no. 16–18, pp. 3818–3829, Oct. 2009.
- [4] N. Wang, "A Generalized Ellipsoidal Basis Function Based Online Self-constructing Fuzzy Neural Network," Springer US, Neural Processing Letters, vol. 34, no. 1, pp. 13–37, Aug. 2011.
- [5] J. D. J. Rubio, "SOFMLS: Online self-organizing fuzzy modified least-squares network," IEEE Trans. Fuzzy Syst., vol. 17, no. 6, pp. 1296–1309, Dec. 2009.
- [6] G. Leng, T. M. McGinnity, and G. Prasad, "Design for self-organizing fuzzy neural networks based on genetic algorithms," IEEE Trans. Fuzzy Syst., vol. 14, no. 6, pp. 755–766, Dec. 2006.
- [7] J. Alcalá-Fdez, R. Alcalá, M. J. Gacto, and F. Herrera, "Learning the membership function contexts for mining fuzzy association rules by using genetic algorithms," Fuzzy Sets Syst., vol. 160, no. 7, pp. 905–921, Jul. 2009.
- [8] Ch. Hung Chen, Ch. J. Lin and Y. Y. Liao, "A Rule-Based Symbiotic Modified Differential Evolution for Self-Organizing Neuro-Fuzzy Systems," Proceedings of International Conference on System Science and Engineering, Macau, China – Jun. 2011.
- [9] Sh. F. Lin, J. W. Chang and Y. C. H., "A self-organization mining based hybrid evolution learning for TSK-type fuzzy model design," Springer US, *Applied Intelligence*, Vol. 36, no. 2, pp. 454–471, Mar. 2012.
- [10] C. F. Juang, S. H. Chiu, and S. W. Chang, "A self-organizing TS-Type fuzzy network with support vector learning and its application to classification problems," IEEE Trans. Fuzzy Syst., vol. 15, no. 5, pp. 998–1008, Oct. 2007.
- [11] C. F. Juang and S. J. Shiu, "Using self-organizing fuzzy network with support vector learning for face detection in color images," Neurocomput., vol. 71, no. 16–18, pp. 3409–3420, Oct. 2008.
- [12] K. Dahal, Kh. Almejalli, M. Alamgir Hossain, W. Chen, "Algorithm-based learning for rule identification in fuzzy neural networks," Appl. Soft Comput. 35, 605–617, 2015.
- [13] N. N. Nguyen, W. J. Zhou, Ch. Quek, "GSETSK: a generic self-evolving TSK fuzzy neural network with a novel Hebbian-based rule reduction approach," Appl. Soft Comput. 35, 92–102, 2015.
- [14] J. Tavoosi, A. A. Suratgar, M. B. Menhaj, "Nonlinear system identification based on a self-organizing type-2 fuzzy RBFN," Eng. Appl. Artif. Intell. 54, 26–38, 2016.
- [15] J. Tavoosi, M. A. Badamchizadeh, "A class of type-2 fuzzy neural networks for nonlinear dynamical system identification," Neural Comput & Applic, Vol 23, no 3, pp. 707–717, Dec. 2013.



روشهایی در تحلیل قابلیت اطمینان سیستمهای تعمیرپذیر در شرایط نادقیق

رضا زارعی *

استادیار گروه آمار، دانشکده علوم ریاضی، دانشگاه گیلان

sheikhi.negar1611@gmail.com

چکیده. یکی از مشخصه‌های مهم تولیدات صنعتی، بحث عملکرد محصول در طول زمان است به طوریکه به نحو مطلوب وظیفه خود را انجام دهد. قابلیت اعتماد آماری مجموعه‌ای از فنون مهندسی و آماری است که به مطالعه خواص تصادفی و برآورد شاخصهای مختلف طول عمر محصولات و سیستمهای مختلف می‌پردازد. با وجود کارآمدی روش‌های موجود در تحلیل قابلیت اطمینان کلاسیک، در عمل با وضعیت‌هایی مواجه می‌شویم که در آن‌ها عدم اطمینان در شرایط و اطلاعات مسئله حکمفرماست. بنابراین ضروری به نظر می‌رسد تا روش‌هایی برای تحلیل قابلیت اطمینان سیستم‌ها در شرایط نادقیق مورد مطالعه قرار گیرد. از طرفی دیگر با توجه به این که خرابی یک سیستم در ارتباط با محصولات امری اجتناب‌ناپذیر است و علاوه بر آن امکان تعویض قطعات معیوب همواره وجود ندارد، تحلیل موارد ذکرشده برای سیستم‌های قابل تعمیر از موضوعات مهم و پرکاربرد است. در این مقاله ضمن معرفی مفاهیم پایه قابلیت اطمینان سیستمهای تعمیرپذیر، برخی از روش‌ها برای تجزیه و تحلیل این سیستم‌ها مورد بررسی قرار گرفته است.

۱. مقدمه

در سال‌های اخیر کاربرد علم آمار در شاخه‌های مختلف علمی سبب شده است که این رشته به عنوان ابزاری قوی در توسعه علوم و فناوری مطرح شود. امروزه به جرات می‌توان ادعا کرد

واژگان کلیدی. حساب اعداد فازی، روش لاند-تاو، سیستم‌های تعمیرپذیر، شبکه پتری، قابلیت اطمینان.

* سخنران

که استفاده از فنون آماری در تحلیل پدیده های گوناگون تصادفی امری اجتناب ناپذیر است. استفاده از علم آمار در مدل سازی و تحلیل پدیده های تصادفی، در شاخه های مختلف مهندسی و صنعت یکی از بسترهای مهم کاربرد این علم است. پیشرفت فناوری و بالا رفتن سطح زندگی مردم با استفاده از دستاوردهای صنعتی، رقابت شدیدی بین تولیدکنندگان برای ارائه محصولات و خدمات مرغوبتر پدیدآورده است. در تولید محصولات راهبردی، چنانچه یک شرکت تولیدی نتواند کیفیت و قابلیت محصولات خود را بر اساس استانداردها مطلوب ارائه کند، دیر یا زود از گردونه رقابت و تولید خارج می شود. توسعه سریع اقتصادی و تجارت آزاد بین المللی سبب شده است که محصولات مشابه از کشورهای متفاوت وارد بازار شده و وجود منابع مختلف اطلاعاتی، به مشتریان این امکان را می دهد که خصوصیات محصولات را از نقطه نظر قیمت، خدمات، قابلیت اعتماد و جنیه های دیگر مقایسه نموده و بهترین گزینه ممکن را انتخاب کنند. با توجه به این موضوع، معمولاً شرکت های معتبر تولیدی در این رقابت به استفاده از فنون مختلف برای بالا بردن قابلیت اعتماد محصولات خود نیاز دارند. در این بین، یکی از مشخصه های مهم تولیدات صنعتی، بحث عملکرد محصول در طول زمان است به طوری که به نحو مطلوب وظیفه خود را انجام دهد. قابلیت اعتماد آماری مجموعه ای از فنون مهندسی و آماری است که به مطالعه خواص تصادفی و برآورد شاخص های مختلف طول عمر محصولات و سیستم های مختلف می پردازد. از طرفی دیگر تمامی شرکت های تولیدی علاقه مند هستند تا از تجربیات و اطلاعات گذشته که در مسیر تولید و ارائه خدمات به دست آورده اند، برای بهبود کیفیت و بالا بردن قابلیت اعتماد محصولات و سیستم های نوین ارائه شده بهره گیرند. با وجود کارآمدی روش های موجود در تحلیل قابلیت اطمینان محصولات، در عمل با وضعیت هایی مواجه می شویم که در آنها عدم اطمینان در شرایط و اطلاعات مسئله حکمفرماست. در چنین مواردی، استفاده از روش های دقیق برای تحلیل قابلیت اطمینان این محصولات می تواند خطایی جبران ناپذیر را به همراه داشته باشد. تا کنون مطالعات و تحقیقات بسیاری در زمینه معرفی روش هایی برای تحلیل قابلیت اطمینان سیستم ها در محیط فازی صورت گرفته است (زارعی ۱۳۸۸). با این

وجود تا جایی که نگارنده می داند، در زمینه تحلیل قابلیت اطمینان سیستم های تعمیرپذیر در شرایط نادقیق، یعنی شرایطی که داده های حاصل از مطالعات، پارامترهای مدل و احتمال های مربوط به سیستم نادقیق بوده اند، مطالعات اندکی صورت گرفته است. تحلیل قابلیت اطمینان سیستم های تعمیرپذیر بر اساس روش های نادقیق برای نخستین بار توسط کنزوچ و اودووم انجام شد **اودووم (۲۰۰۱)**. در تحقیق انجام شده توسط آن ها از روش لاند-تاو و شبکه های پتری برای این منظور استفاده شد. پس از آن برخی مطالعات دیگر نیز صورت گرفته است که مهم ترین آن ها به کوتاهی مرور میشوند. گارگ ؟ با توسعه رویکرد ارائه شده توسط کنزوچ و اودووم، تحلیل قابلیت اطمینان سیستم های تعمیرپذیر را بر اساس مجموعه های مبهم پیشنهاد نمود. کومال و شارما **کومال (۲۰۱۴)** با ترکیب روش لاند-تاو با شبکه های عصبی مصنوعی رویکردی جدید را برای تحلیل قابلیت اطمینان این نوع سیستم ها ارائه نمودند. گارگ و همکاران با ترکیب الگوریتم کلونی زنبور عسل با روش لاند-تاو توانستند روشی توانمند و جدید را برای ارزیابی قابلیت اطمینان سیستم های تعمیرپذیر پیشنهاد دهند **گارگ و همکاران (۲۰۱۴)**. در این مقاله ضمن معرفی مفاهیم مربوط به قابلیت اطمینان سیستم های تعمیرپذیر، روشی برای تحلیل این سیستم ها در محیط فازی مورد بررسی قرار خواهد گرفت. در ادامه و در بخش دوم، مفاهیم و مقدماتی در رابطه با نظریه قابلیت اطمینان کلاسیک و برخی روش های تحلیل آن ارائه شده است. در بخش سوم، تجزیه و تحلیل قابلیت اطمینان سیستم های تعمیرپذیر مورد بحث و بررسی قرار گرفته است. در پایان و در بخش چهارم، ضمن بیان مقدماتی کوتاه از اعداد فازی، به تحلیل و ارزیابی قابلیت اطمینان سیستم های تعمیرپذیر تحت شرایط نادقیق پرداخته شده است.

۲. تعاریف و مفاهیم مورد نیاز

در این بخش برخی از تعاریف و مفاهیم مورد نیاز مرتبط با نظریه قابلیت اطمینان تا آنجایی که در بخش های بعد به آن ها نیازمند هستیم را مرور می کنیم.

۱.۲. قابلیت اطمینان سیستم و برخی مفاهیم سالخوردگی. فرض کنید یک سیستم (منظور از سیستم میتواند یک قطعه الکتریکی و یا هر وسیله یا شی دیگر باشد) دارای طول عمری تصادفی باشد که بتوان آن را توسط متغیر تصادفی T مدل بندی کرد. در این صورت قابلیت اطمینان سیستم به این معنی است که سیستمی با طول عمر T بتواند وظایف محول شده را در طی یک دوره زمانی مشخص به درستی انجام دهد. به عبارت دقیق تر، تابع قابلیت اطمینان در لحظه t که آن را با $R(t)$ نمادگذاری می کنیم، به صورت زیر تعریف می شود

$$R(t) = \int_t^{\infty} f(x)dx,$$

که در آن $f(\cdot)$ تابع چگالی احتمال متغیر تصادفی T است (برای حالت گسسته تعریف مشابه وجود دارد). یکی از مهم ترین معیارها در مطالعات قابلیت اطمینان سیستم ها و مباحث مرتبط با آن، تابع نرخ خرابی می باشد که به این سوال اساسی پاسخ می دهد؛ با اطلاع از اینکه میدانیم سیستم تحت بررسی کماکان در لحظه t هنوز فعال است، چقدر احتمال دارد که بلافاصله پس از زمان t از کار بیافتد. در این شرایط تابع نرخ خرابی در لحظه t که آنرا با $\lambda(t)$ نشان می دهیم از رابطه زیر به دست می آید [۴]

$$\lambda(t) = \frac{f(t)}{R(t)}.$$

ازجمله مباحثی که در نظریه قابلیت اطمینان از اهمیت ویژه ای برخوردار است، تحلیل قابلیت اطمینان سیستم هایی است که در صورت از کار افتادگی می توانند در فرایندی مشخص دوباره به وضعیت فعالیت اولیه برگردند. بیشتر سیستم ها و تجهیزات پیشرفته که امروزه در زندگی بشر مورد استفاده قرار می گیرند از این نوع هستند که آن ها را سیستم های تعمیرپذیر می گویند. هواپیما، اتومبیل، کامپیوتر و ... از نوع سیستم های تعمیرپذیر هستند. بدیهی است که چنین سیستم هایی دارای قطعاتی اصلی و کلیدی هستند که معمولاً قیمت بالایی داشته و انتظار می رود که طول عمر زیادی داشته باشند. از اینرو، در صورت بروز هرگونه نقص یا عیب احتمالی، مشتری ترجیح می دهد با توجه به هزینه بالای قطعات، آن ها را تعمیر کند تا تعویض. بنابراین،

در سیستم های تعمیرپذیر، علاوه بر نرخ خرابی مفهومی تحت عنوان نرخ تعمیر سیستم نیز مورد توجه قرار می گیرد که معمولاً آن را با τ نشان می دهند. ^[۴]

از دیدگاهی دیگر، قابلیت اطمینان یک سیستم بر اساس ساختار شکل گیری اجزاء آن مورد بررسی قرار میگیرد که در اصطلاح قابلیت اطمینان ساختاری سیستم نامیده می شود. در این رویکرد سیستم شامل تعدادی معین جزء است و برای اجرای هدفی معین طراحی شده است. بدیهی است که عملکرد کل سیستم تابعی از عملکرد اجزای آن باشد. با فرض اینکه هر جزء از سیستم تنها می تواند در دو وضعیت کاملاً فعال و کاملاً غیرفعال باشد، وضعیت کل سیستم نیز در یکی از دو وضعیت کاملاً فعال و یا کاملاً غیرفعال خواهد بود. برای تعیین وضعیت سیستم بر حسب وضعیت اجزای آن، فرض می شود که رابطه سیستم با اجزای آن با تابع دودویی $\phi(\underline{x})$ که به تابع ساختار سیستم معروف است، تعیین می گردد که در آن بردار $\underline{x}$ را بردار وضعیت اجزای سیستم گویند. بر این اساس تابع ساختار سیستم به صورت زیر تعریف می شود

$$\phi(\underline{x}) = \begin{cases} 1 & \text{اگر سیستم فعال باشد} \\ 0 & \text{اگر سیستم غیرفعال باشد} \end{cases}$$

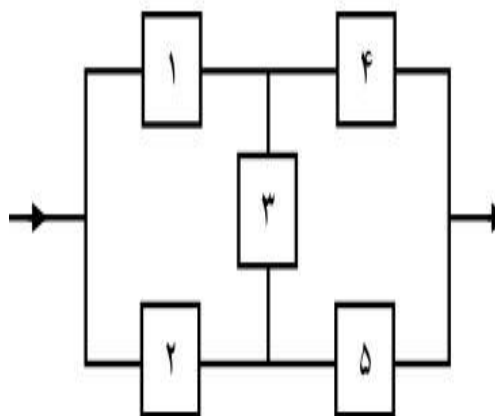
در این رویکرد قابلیت اطمینان سیستم از رابطه زیر به دست می آید

$$R = P(\phi(\underline{x}))$$

که در آن احتمال مربوطه بستگی به وضعیت قرار گرفتن اجزاء در سیستم مورد بحث قرار گرفته و تعیین می شود.

مثال ۱۰۲. فرض کنید سیستمی دارای پنج جزء باشد که نحوه قرار گرفتن آن ها در شکل ۱ نمایش داده شده است

با توجه به این که هر جزء تنها می تواند در دو وضعیت ممکن قرار بگیرد، بنابراین تعداد سی و دو بردار مسیر ممکن وجود دارد که در شانزده مورد تابع ساختار برابر با یک خواهد بود و

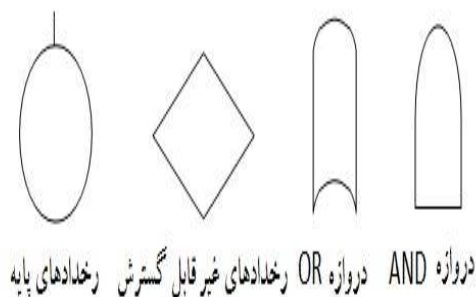


شکل ۱: نحوه قرار گرفتن اجزاء در مثال ۱.۲

قابلیت اطمینان سیستم برحسب این بردارها تعیین خواهد شد. مجموعه های مسیر عبور و قطع کننده مینیمال یکی از روش های کلیدی برای به دست آوردن قابلیت اطمینان ساختاری سیستم ها استفاده از مفهوم مجموعه مسیرهای عبور و قطع کننده مینیمال است. در تعریف، یک مسیر مینیمال بردار وضعیتی است که به ازای آن سیستم فعال است و اگر حداقل یکی از اجزای فعال آن بردار، غیرفعال شود، سیستم نیز غیرفعال خواهد شد. مشابه با آن، بردار قطع کننده مینیمال برداری است که به ازای آن سیستم از کار می افتد و چنان چه حداقل یکی از اجزای آن بردار فعال شود، ممکن است سیستم شروع به کار کند. مزیت استفاده از این روش این است که تعداد زیاد مسیرهای ممکن توسط مسیرهای عبور و قطع کننده مینیمال که تعداد آن ها با توجه به تعریف کمتر خواهد بود، جایگزین شده و از این رو حجم محاسبات به شکل چشمگیری کاهش خواهد یافت. مطابق با مفاهیم ارائه شده برای مجموعه های عبور و قطع کننده مینیمال، در مثال ۱.۲ مجموعه $(1, 4)$ ، $(2, 5)$ ، $(1, 3, 5)$ و $(2, 3, 4)$ مجموعه مسیرهای مینیمال و مجموعه قطع کننده های مینیمال به صورت $(1, 2)$ ، $(4, 5)$ ، $(1, 3, 5)$ و $(2, 3, 4)$ می باشند. حال با داشتن این مجموعه های مینیمال می توان قابلیت اطمینان سیستم را به دست آورد.

۳. تحلیل قابلیت اطمینان سیستم ها بر اساس تحلیل درخت عیب

تحلیل درخت عیب روشی گرافیکی برای تعیین قابلیت اطمینان سیستم های تحت بررسی است که در آن از ریشه یابی سطح به سطح عوامل مختلفی که ممکن است سبب خرابی سیستم شوند، استفاده می شود. این عیب ها معمولا با اشتباهات انسانی و خرابی سخت افزاری و یا هر چیزی که مربوط به رخدادهای نامطلوب باشد همراه است. از دیدگاه کلی، تحلیل درخت عیب را می توان یک مدل کیفی دانست که به وسیله این مدل می توان اغلب مدل های کمی را برآورد کرد. به دلیل پیچیده بودن ساختار و عملکرد تحلیل درخت عیب برای راحتی و فهم مطلب از دروازه هایی استفاده می کنند (شکل ۲ را ببینید) که هر کدام از این دروازه ها با ورودی و خروجی همراه هستند. دروازه های درخت عیب بر دو نوع OR و AND می باشند. از دروازه OR زمانی



شکل ۲: برخی نمادهای درخت عیب

استفاده می کنیم که حداقل یکی از ورودی های اتفاق افتاده باشند. خروجی دروازه AND زمانی رخ می دهد که همه ی ورودی ها اتفاق افتاده باشند. در دروازه AND برخلاف دروازه OR یک رابطه علی بین ورودی ها و خروجی عیب ها وجود دارد. حال، با فرض این که احتمال خرابی

ها در سطوح مختلف برای محقق مشخص شده باشد، قادر خواهیم بود تا با استفاده از قوانین احتمال و دروازه های تعریف شده، قابلیت اطمینان کل سیستم را تعیین کنیم.

۴. قابلیت اطمینان سیستم های تعمیرپذیر

در این بخش مفاهیم مقدماتی مربوط به تحلیل قابلیت اطمینان سیستم های تعمیرپذیر مورد بررسی قرار می گیرد. در این راستا، ابتدا از تکنیکی موسوم به تکنیک لاند-تاو استفاده شده و سپس برای تسهیل استفاده از آن به معرفی شبکه های پتری خواهیم پرداخت.

۱.۴. تکنیک لاند-تاو در تحلیل قابلیت اطمینان سیستم های تعمیرپذیر. یک سیستم تعمیرپذیر شامل دو جزء را به ترتیب با نرخ خرابی و نرخ تعمیر λ و τ در نظر گرفته و فرض کنید که λ_1 و λ_2 نرخ شکست و τ_1 و τ_2 نرخ تعمیر اجزاء سیستم باشند. علاوه بر آن، فرض کنید $\lambda_1, \lambda_2, \tau_1, \tau_2$ مقادیری ثابت و λ, τ مقادیر نسبتاً کوچکی باشند تحت این فرض ها، پیشامد E را با تعریف زیر در نظر بگیرید (شکل ۳ را ملاحظه کنید)

حال فرض کنید A پیشامدی از است که برای اولین بار در فاصله dt اتفاق می افتد و B پیشامدی از است که بین t-، t اتفاق می افتد. اگرچه A و B مستقل هستند، اما $rA =$ بنابرین می توان نوشت $PrB = F(t) - F(t-)$

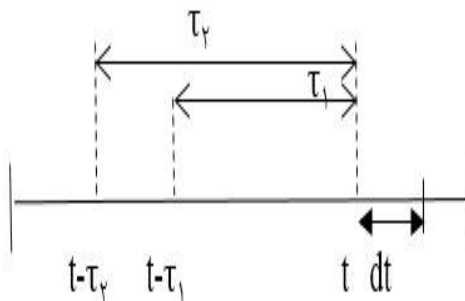
$$Pr\{B\} = \{1 - e^{-\lambda_2 t}\} - \{1 - e^{-\lambda_2(t-\tau_2)}\} = -e^{-\lambda_2 t} \{1 - e^{\lambda_2 \tau_2}\}$$

در این صورت

$$dp_1 = Pr\{A\} \times Pr\{B\} = -\lambda_1 e^{-\lambda_1 t} \{1 - e^{-\lambda_2 \tau_2}\} dt$$

بنابراین

$$p_1 = \frac{\lambda_1}{\lambda_2} (e^{\lambda_2 \tau_2} - 1) (e^{-\lambda_1 T} - 1)$$



شکل ۳: نمودار تحلیلی سیستم تعمیرپذیر دارای دو جزء

حال اگر مقادیر λT و $\lambda \tau$ خیلی کوچک باشند، خواهیم داشت

$$P_1 \approx \frac{\lambda_1 (\lambda_2 \tau_2) (\lambda_2 T)}{\lambda_2} = \lambda_1 \lambda_2 \tau_2 T$$

به طور مشابه $P_2 = \lambda_2 \lambda_1 \tau_1 T$ به دست می آید. بنابراین

$$\Pr \{E\} = \lambda_1 \lambda_2 \lambda_3 (\tau_2 \tau_3 + \tau_1 \tau_3 + \tau_1 \tau_2) T$$

در حالتی که سیستم تحت بررسی دارای m جزء باشد، با تعمیم مراحل بالا می توان نشان داد

$$\Pr \{E\} = \prod_{j=1}^m \lambda_j \left[\sum_{i=1}^m \prod_{j=1, i \neq 1}^m \tau_j \right] T. \quad (1.4)$$

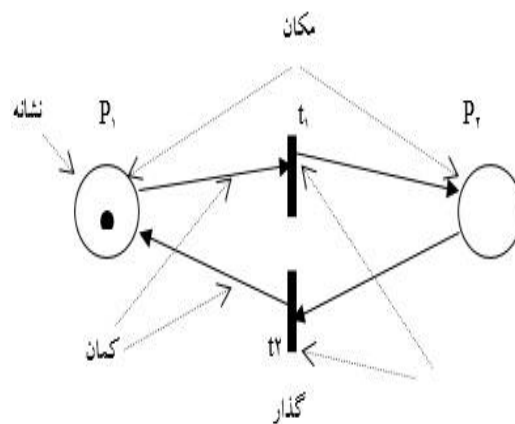
از طرفی می دانیم در دروازه AND باید تمامی شکست ها هم زمان اتفاق بیفتند و در دروازه OR کافی است حداقل یکی از شکست ها ممکن اتفاق بیفتد. بنابراین نرخ خرابی و نرخ تعمیر در دروازه AND به صورت زیر به دست می آید

$$\lambda_{AND} = \prod_{j=1}^m \lambda_j \left[\sum_{i=1}^m \prod_{j=1, i \neq 1}^m \tau_j \right], \quad \tau_{AND} = \frac{1}{\sum_{j=1}^m \frac{1}{\tau_j}}. \quad (2.4)$$

به طور مشابه نرخ خرابی و نرخ تعمیر برای دروازه OR به صورت زیر حاصل می شوند

$$\lambda_{OR} = \sum_{j=1}^m \lambda_j, \quad \tau_{OR} = \frac{\sum_{j=1}^m \lambda_j \tau_j}{\sum_{j=1}^m \lambda_j}. \quad (3.4)$$

۲.۴. شبکه های پتری. نظریه شبکه های پتری برای نخستین بار توسط آدام پتری در سال ۱۹۶۲ ارائه شد پتری (۱۹۶۲). این شبکه ها یک زبان مدل سازی ریاضیاتی برای توصیف سیستم های توزیع شده می باشد. ساختار شبکه پتری از چهار جزء مکان، گذار، نشانه و کمان تشکیل شده است که در شکل ۴ نمایش داده شده اند.



شکل ۴: شمای کلی یک شبکه پتری

علاوه بر آن، شبکه های پتری از مجموعه قواعد زیر پیروی می کنند:

الف) یک کمان نمی تواند دو مکان و یا دو گذار را به هم وصل کند.

ب) یک کمان می تواند به چند گذار وصل شود و چند مکان به یک گذار و بر عکس این هم وجود دارد. پ) نشانه در مکان قرار می گیرد نه گذار.

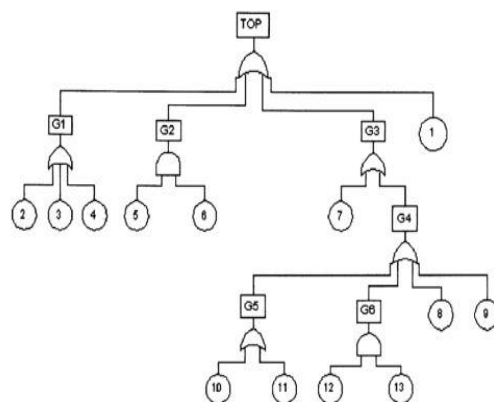
ت) یک گذار تواناست اگر هر مکان ورودی آن حداقل یک نشانه داشته باشد.

ث) شلیک شدن به صورت اتوماتیک و پس از فعال شدن گذار انجام می شود.

ج) وقتی يك گذر شليك مي گردد تمام مهره ها فعال کننده را از مكان هاي ورودی خود خارج و سپس به مكان هاي خروجی، ارسال مي کند.

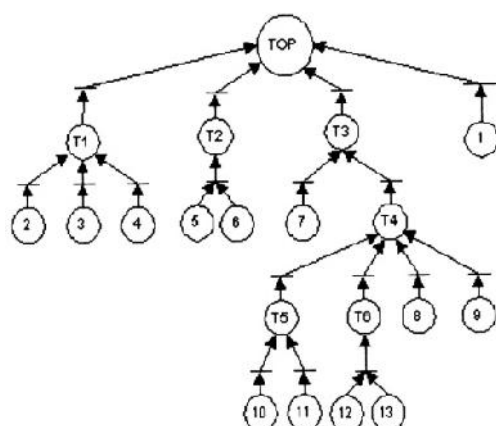
از بین مزیت های استفاده از شبکه پتری به جای تحلیل درخت عیب به دو دلیل عمده بسنده می کنیم. نخست آن که برای تحلیل قابلیت اطمینان بر اساس درخت عیب یک شبکه پتری معادل وجود دارد که بر اساس قوانین و سازوکار شبکه های پتری قابل ترسیم است. اما در مدل بندی بر اساس درخت عیب، محقق ناچار است که از نمادهای مختلفی برای پیشامدهای خرابی گوناگونی که ممکن است رخداد کنند، استفاده کند. این در حالیست که با توجه به ساختار شبکه پتری و قوانین حاکم بر آن این مشکل برطرف خواهد شد.

برای مثال فرض کنید که درخت عیب یک سیستم تحت بررسی پس از ریشه یابی رخدادهای نامطلوب ممکن و سطح بندی آن، مدل بندی شده و نمودار آن در شکل ۵ ارائه شده است. در این صورت می توان شبکه پتری معادل با این درخت عیب را همان طور که در شکل ۶ نمایش داده شده است، صورت بندی نمود.



شکل ۵: درخت عیب سیستم مفروض

اکنون که می توان روش کارکرد سیستم را در قالب یک شبکه پتری نمایش داد، مزیت دوم امکان استفاده از مجموعه مسیرهای عبور و قطع کننده مینیمال است. از بین روش های مختلفی که برای



شکل ۶: شبکه پتری معادل با درخت عیب سیستم مفروض

بکارگیری این تکنیک تاکنون معرفی شده اند، استفاده از روش ماتریسی به نظر مرسوم تر و البته ساده تر است **لیو (۱۹۹۷)**. بر اساس این روش، در سیستم مورد بحث که نمایش شبکه پتری آن در شکل ۶ آورده شد، مجموعه های مسیر عبور مینیمال شامل (۱،۲،۳،۴،۵،۷،۸،۹،۱۰،۱۱،۱۲) و (۱،۲،۳،۴،۶،۷،۸،۹،۱۰،۱۱،۱۲) و مجموعه های قطع کننده مینیمال شامل (۱)، (۲)، (۳)، (۴)، (۵،۶)، (۷)، (۸)، (۹)، (۱۰)، (۱۱) و (۱۲،۱۳) می باشند. به وضوح می توان دید که حجم محاسبات در این حالت بسیار کم شده است. این موضوع به ویژه در تحلیل سیستم های پیچیده تر که معمولا در شرایط واقعی با آن روبرو خواهیم شد بسیار راهگشا خواهد بود.

۵. قابلیت اطمینان سیستم های تعمیرپذیر در شرایط نادقیق

در تحلیل قابلیت اطمینان سیستم های معمولی و تعمیرپذیر فرض های اساسی دودویی بودن اجزاء و توصیف رفتارشان در قالب الگوهای احتمالی همواره بایستی برقرار باشند. با این وجود، در شرایط واقعی به سبب شرایط حاکم بر مسئله، اشتباهات انسان یا ماشین و نیز ماهیت داده های جمع آوری شده از سیستم تحت بررسی قادر نخواهیم بود تا داده های حاصل از مطالعات، پارامترهای مدل و احتمال های مربوط به کارکرد اجزاء سیستم را به طور دقیق اندازه گیری کنیم.

قدرمسلّم در چنین وضعیت هایی استفاده از روش های کلاسیک جهت ارزیابی قابلیت اطمینان این سیستم ها، نتایجی نادرست را به دنبال خواهد داشت. در مقابل، محقق می تواند با بهره گیری از ابزار توانمندی که نظریه مجموعه های فازی جهت صورت بندی مفاهیم نادقیق در اختیار وی قرار داده است، روشی کارا را برای مدل بندی قابلیت اطمینان چنین سیستم هایی ارائه کند. اعداد فازی از مهم ترین این مفاهیم می باشد که با استفاده از آن می توان مفاهیم نادقیق موجود در مسئله را صورت بندی نمود. بر اساس تعاریف موجود، عدد فازی را یک عدد فازی مثلثی گوئیم هرگاه تابع عضویت آن به صورت زیر تعریف شده باشد

$$\mu_{\tilde{A}}(x) = \begin{cases} \frac{x-a_1}{a_2-a_1} & a_1 \leq x \leq a_2 \\ \frac{a_3-x}{a_3-a_2} & a_2 \leq x \leq a_3 \\ 0 & o.w \end{cases}$$

که در آن، و اعدادی حقیقی می باشند. علاوه بر آن، α -برش این عدد فازی مثلثی بر اساس تعریف از رابطه زیر به دست می آید

$$\tilde{A}_\alpha = [(a_2 - a_1)\alpha + a_1, (a_3 - a_2)\alpha + a_3] \quad (1.5)$$

پس از معرفی اعداد فازی، تحقیقات بسیاری برای تعریف عملگرهای مجموعه ای بین این اعداد صورت گرفته است که دو رویکرد اصلی استفاده از اصل گسترش و تعریف این عملگرها بر اساس برش اعداد فازی می باشد. بر این اساس، در صورتی که و دو عدد فازی به ترتیب با α -برش های $\tilde{A}_\alpha = [\tilde{A}_\alpha^L, \tilde{A}_\alpha^U]$ و $\tilde{B}_\alpha = [\tilde{B}_\alpha^L, \tilde{B}_\alpha^U]$ باشند، آن گاه برای مثال جمع این اعداد فازی از رابطه زیر به دست می آید

$$\tilde{A} + \tilde{B} = \tilde{A}_\alpha + \tilde{B}_\alpha = [\tilde{A}_\alpha^L + \tilde{B}_\alpha^L, \tilde{A}_\alpha^U + \tilde{B}_\alpha^U]. \quad (2.5)$$

اکنون با استفاده از این مفاهیم و مطالب بخش های پیشین سعی داریم تا به تحلیل قابلیت اطمینان سیستم های تعمیرپذیر بپردازیم.

با در نظر گرفتن شرایط نادقیقی که به آن ها اشاره شد، فرض کنید که نرخ خرابی و نرخ تعمیر سیستم تحت بررسی را نمی توان به طور دقیق ثبت نمود و این دو کمیت را تنها می توان در قالب عبارات زبانی مانند تقریبا و حدودا بیان نمود. در این وضعیت ها، می توان با استفاده از اطلاعات گردآوری شده از مسئله، اطلاعات پیشین و یا نظر خبره این مفاهیم را در قالب اعداد فازی و منعکس نمود. با در نظر گرفتن فرم مثلثی برای این اعداد فازی و استفاده از روابط ۱.۵ و ۲.۵، برش های نرخ خطر فازی و نرخ تعمیر فازی سیستم را در دروازه AND می توان به صورت زیر به دست آورد

$$(3.5) \quad \tilde{\lambda}_{\alpha} = \left[\prod_{i=1}^n ((\lambda_{i2} - \lambda_{i1}) \alpha + \lambda_{i1}) \cdot \sum_{j=1}^n \left[\prod_{\substack{i=1 \\ j \neq i}}^n ((\tau_{i2} - \tau_{i1}) \alpha + \tau_{i1}) \right], \right. \\ \left. \prod_{i=1}^n (-(\lambda_{i3} - \lambda_{i2}) \alpha + \lambda_{i3}) \cdot \sum_{j=1}^n \left[\prod_{\substack{i=1 \\ j \neq i}}^n (-(\tau_{i3} - \tau_{i2}) \alpha + \tau_{i3}) \right] \right]$$

به طور مشابه می توان روابط زیر را برای α -برش های نرخ خطر فازی و نرخ تعمیر فازی سیستم را در دروازه OR به صورت زیر به دست آورد

$$(4.5) \quad \tilde{\tau}_{\alpha} = \left[\frac{\prod_{i=1}^n ((\tau_{i2} - \tau_{i1}) \alpha + \tau_{i1})}{\sum_{j=1}^n \left[\prod_{\substack{i=1 \\ j \neq i}}^n (-(\tau_{i3} - \tau_{i2}) \alpha + \tau_{i3}) \right]}, \frac{\prod_{i=1}^n (-(\tau_{i3} - \tau_{i2}) \alpha + \tau_{i3})}{\sum_{j=1}^n \left[\prod_{\substack{i=1 \\ j \neq i}}^n (-(\tau_{i3} - \tau_{i2}) \alpha + \tau_{i3}) \right]} \right]$$

با داشتن α -برش های نرخ خرابی و نرخ تعمیر فازی سیستم و با استفاده از اتحاد تجزیه قادر هستیم تا تابع عضویت مربوط به این دو کمیت فازی را به دست آورده و بر اساس آن به تجزیه و تحلیل وضعیت سالخوردگی سیستم در سطوح مختلف بپردازیم.

مثال ۱.۵. واحد خمیرسازی کاغذ شامل بخش های مختلفی مانند پخت خرده های چوب، جداسازی گره ها توسط ماشین گره زنی، شستشو خمیر کاغذ با استفاده از لیکور سفید و باز کردن الیاف توسط

بازکننده‌های الیاف می‌باشد. ابتدا خرده‌های چوب برای پخته شدن در انباری ذخیره می‌گردند که به دیجستر موسوم است. بعد از پخت خرده‌های چوب و مخلوط شدن آن‌ها با یکدیگر، سود را به آن اضافه کرده و سپس پخت آن همراه با بخار خشک طی چند ساعت در دیگ پخت خمیر، انجام می‌گیرد. خرده چوب‌های پخته‌شده دارای چندین گره می‌باشد که این گره‌ها مانع تولید کاغذ می‌شوند. در این مرحله گره‌های ایجادشده توسط ماشین گره‌زنی باز می‌شوند. بعد از باز کردن گره‌ها، این خمیر به ظرف استوانه شکل که به دکر معروف است انتقال می‌یابند. وظیفه دکر از بین بردن لیکور سیاه موجود در خمیر کاغذ تا حد امکان می‌باشد و سپس شستشو در چندین مرحله دیگر نیز ادامه می‌یابد. سرانجام در آخرین مرحله برای جداسازی الیاف خمیر از یکدیگر، از دستگاه بازکننده الیاف که با چرخش و سرعت بالا انجام می‌گیرد، استفاده می‌شود و خمیر کاغذ به صورت آماده و با الیاف بسیار ریز درمی‌آیند.

واحد خمیر سازی کاغذ شامل زیرسیستم‌های زیر می‌باشد:

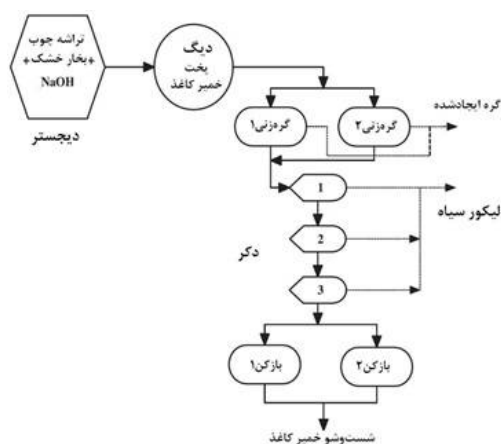
دیجستر: دستگاهی برای پخت خرده‌های چوب می‌باشد. شکست این زیرسیستم باعث شکست کل سیستم می‌شود.

ماشین گره‌زنی: این بخش از سیستم شامل دو زیرسیستم دیگر می‌باشد. زیرسیستم اول، سیستم در حال کار کردن باشد و دیگری سیستم در حالت آماده‌به‌کار. ماشین گره‌زنی برای از بین بردن گره‌ها در خرده‌های چوب که مانع تولید کاغذ می‌شوند به کار گرفته می‌شود. سیستم خمیرسازی کاغذ زمانی با شکست مواجه می‌شود که هر دو زیرسیستم این زیرسیستم اصلی با شکست روبرو شود.

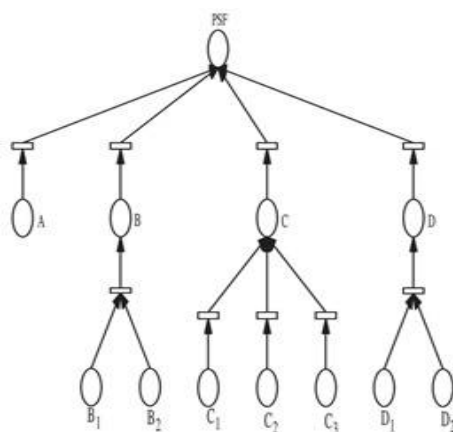
دکر: دکر شامل سه بخش است که به صورت سری بسته می‌شوند. وظیفه این زیرسیستم از بین بردن لیکور سیاه موجود در خمیر کاغذ می‌باشد. شکست هرکدام از این سه بخش موجب شکست کل سیستم خمیرسازی کاغذ خواهد شد.

بازکننده الیاف: این زیرسیستم نیز به مانند ماشین گره‌زنی شامل دو زیرسیستم در حال کار و آماده به کار می‌باشد و الیاف‌ها را از هم جدا می‌کند. اگر هر دو زیرسیستم این زیرسیستم اصلی با

شکست مواجه شود آنگاه، کل سیستم خمیر سازی کاغذ با شکست مواجه خواهد شد. روش کارکرد واحد خمیر سازی کاغذ همراه با مدل شبکه پتری متناظر آن به ترتیب در شکل ۷ و شکل ۸ آورده شده است که در آن عبارت *PFS* نشان‌دهنده شکست سیستم خمیر سازی کاغذ می‌باشد.



شکل ۷: تشریح روش کارکرد سیستم خمیر سازی در مثال ۱.۵



شکل ۸: شبکه پتری متناظر با سیستم خمیر سازی در مثال ۱.۵

برای تحلیل وضعیت سیستم تحت بررسی، ابتدا اطلاعات را از منابع مختلف جمع آوری نموده و سپس بر اساس نظر کارشناسان خبره کامل تر می شود. داده های مربوط به اجزای زیرسیستم ها در جدول ۱ آورده شده است. حال با در نظر گرفتن این نکته که اطلاعات گردآوری شده برای هر یک از زیرسیستم ها تحت شرایط و عملکردهایی متفاوت به دست آمده که هر کدام نیز با خطای انسانی و ماشینی در ثبت اطلاعات همراه بوده اند، شرایط نادقیقی را برای این مجموعه اطلاعات در نظر می گیریم.

برای این منظور، بر اساس نظر کارشناس خبره هر مجموعه از اطلاعات با در نظر گرفتن عدم اطمینان مناسب موجود و تعیین آن، داده ها را در قالب اعداد فازی مثلثی مناسب صورت بندی و ارائه می کنیم. سپس با استفاده از روابط ۲.۵ و ۴.۵، α -برش های نرخ خطر فازی و نرخ تعمیر فازی سیستم را به دست می آوریم. در مرحله بعدی و به منظور تحلیل و ارزیابی قابلیت اطمینان کل سیستم، از مدل شبکه پتری متناظر آن استفاده می کنیم.

در این مثال، مجموعه های قطع کننده مینیمال سیستم عبارتند از

$$(A), (B \cup B_1), (C_1), (C_2), (C_3), (D_1 \cup D_2).$$

در این صورت، بر اساس روش تشریح شده می توانیم به تجزیه و تحلیل قابلیت اطمینان سیستم تحت بررسی بپردازیم. برای مثال، فرض کنید محقق برای صورت بندی شرایط عدم اطمینان موجود در اطلاعات جمع آوری شده از پهنه های ۱۵/۰ برای هر مجموعه از اطلاعات گزارش شده در جدول (۱) استفاده کند. بر این اساس و پس از انجام محاسبات عددی، نرخ خرابی فازی برابر با ۶/۶۶۰۳۲ و نرخ تعمیر فازی برابر با ۳/۳۱۵۶۱ به دست خواهند آمد. با در نظر گرفتن پهنه ۲۵/۰، این مقادیر به ترتیب ۶/۶۹۶۴۸ و ۳/۷۶۲۰۴ خواهند بود. این در حالیست که استفاده از روش دقیق در هر دو حالت مقادیر یکسانی را برای این دو معیار ارزیابی سیستم گزارش نموده است. در واقع، بدون در نظر گرفتن شرایط نادقیق در گردآوری اطلاعات، در هر سطحی نرخ خرابی و نرخ تعمیر سیستم به ترتیب برابر با ۶/۶۴۰۰ و ۳/۰۹۰۳۶ خواهد بود.

جدول ۱: اطلاعات مربوط به زیر سیستم های واحد خمیرسازی

دایجستر	گره زنی	دکر	بازکن	
۰/۰۰۰۰۰۴	۰/۰۰۰۵	۰/۰۰۰۲	۰/۰۰۰۵	نرخ خطر
۱۸	۶	۳	۶	نرخ تعمیر

مراجع

- زارعی، ر. (۱۳۸۸). قابلیت اطمینان در محیط فازی، *مجله اندیشه آماری*، ۱۲، ۸۶-۹۹.
- Knezevic, J. and Odoom, E. R. (2001), Reliability modeling of repairable systems using Petri nets and fuzzy Lambda-Tau methodology. *Reliability Engineering and System Safety*, 73(1), 1-17.
- Garg, H. (2013), Reliability analysis of repairable systems using Petri nets and Vague Lambda-Tau methodology, *ISA Transactions*, 52(1), 6-18.
- Garg, H. (2014), Performance and behavior analysis of repairable industrial systems using Vague Lambda-Tau methodology, *Applied Soft Computing*, 22, 323-338.
- Komal, Sharma, S.P. (2014), Fuzzy reliability analysis of repairable industrial systems using soft-computing based hybridized techniques, *Applied Soft Computing*, 24, 264-276.
- Garg, H., Rani, M. and Sharma, S.P. (2014), An approach for analyzing the reliability of industrial systems using soft-computing based technique, *Expert Systems with Applications*, 41, 489-501.
- C.A. Petri, Communication with automata, doctoral thesis, Ph.D. thesis, University of Bonn, Technical Report (English) RADC-TR-65-377,

Rome Air Development Center, Giriffis, NY, 1966, 1962.

Liu, T.S. and Chiou, S.B. (1997). The application of Petri nets to failure analysis, *Reliability Engineering System Safety*, 57, 129-142.

برآورد ريج در مدل رگرسیون فازی با متغیرهای مستقل و پاسخ فازی

سیده راضیه سجادی^۱ *، محسن عارفی^۲، و محمد خراشادی زاده^۳

^۱ دانشگاه بیرجند، دانشکده علوم ریاضی و آمار، گروه آمار

s.r.sajjadi67@gmail.com

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

arefi@birjand.ac.ir

^۳ دانشگاه بیرجند، دانشکده علوم ریاضی و آمار، گروه آمار

m.khorashadizadeh@birjand.ac.ir

چکیده. در این مقاله مدل رگرسیونی خطی چندگانه، با متغیرهای مستقل و پاسخ فازی در حالتی مورد بررسی قرار گرفته است که بین متغیرهای ورودی هم خطی چندگانه وجود داشته باشد. در این حالت برای برخورد با مشکل هم خطی، از رگرسیون ريج، در برآورد پارامترهای مدل استفاده شده است.

۱. پیش‌گفتار

یکی از روش‌های برآورد پارامترهای مدل رگرسیونی، روش کمترین مربعات خطاست. اگر بین یک یا چند متغیر مستقل یک رابطه خطی وجود داشته باشد هم خطی چندگانه رخ می‌دهد که در این حالت معکوس ماتریس $X'X$ وجود ندارد و نمی‌توان از روش کمترین مربعات خطا استفاده کرد. در حالت وجود هم خطی برای برآورد پارامترها از روش‌هایی مانند رگرسیون مولفه‌های اصلی، رگرسیون کمترین مربعات جزئی، رگرسیون ريج و... استفاده می‌شود. در این مقاله برای برآورد پارامترهای مدل خطی چندگانه فازی در زمان وجود هم خطی از روش رگرسیون ريج استفاده شده است.

تعریف ۱.۱. مجموعه فازی N از R (اعداد حقیقی) را یک عدد فازی گوئیم، اگر

(۱) مجموعه N نرمال و تک‌نمایی باشد. یعنی یک و دقیقاً یک $x_0 \in R$ وجود داشته باشد که $N(x_0) = 1$.

(۲) به ازای هر $\alpha \in (0, 1]$ ، α -برش‌های N به صورت بازه‌های بسته و کراندار باشند.

مجموعه همه اعداد فازی را با $F(R)$ نشان می‌دهند.

واژگان کلیدی. رگرسیون ريج، رگرسیون فازی، اعداد فازی.
* سخنران

تعریف ۲.۱. عدد فازی $\tilde{A} = (a, s)_T$ را یک عدد فازی مثلثی متقارن می نامند اگر تابع عضویت آن به صورت زیر باشد:

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

تعریف ۳.۱. فرض کنید $\tilde{A}_1 = (a_1, s_1)_T$ و $\tilde{A}_2 = (a_2, s_2)_T$ دو عدد فازی مثلثی متقارن و c یک عدد حقیقی باشند:

$$c\tilde{A}_1 = (ca_1, |c|s_1)_T \quad (۱)$$

$$\tilde{A}_1 \oplus \tilde{A}_2 = (a_1 + a_2, s_1 + s_2)_T \quad (۲)$$

تعریف ۴.۱. فاصله بین دو عدد فازی $\tilde{A}$ و $\tilde{B}$ بر پایه ی تابع وزنی $f(\alpha)$ به صورت زیر تعریف می شود: [۴]

$$d(\tilde{A}, \tilde{B}) = \left[\int_0^1 f(\alpha) d^2(\tilde{A}_\alpha, \tilde{B}_\alpha) d\alpha \right]^{\frac{1}{2}}$$

که

$$d^2(\tilde{A}_\alpha, \tilde{B}_\alpha) = [a_1(\alpha) - b_1(\alpha)]^2 + [a_2(\alpha) - b_2(\alpha)]^2$$

و $f(\alpha)$ یک تابع صعودی روی بازه $[0, 1]$ است که برای آن $f(0) = 0$ و $\int_0^1 f(\alpha) d\alpha = \frac{1}{2}$.

توجه ۵.۱. اگر $\tilde{A} = (a, s_a)_T$ و $\tilde{B} = (b, s_b)_T$ دو عدد فازی متقارن باشند این فاصله بر پایه تابع وزنی $f(\alpha) = \alpha$ به صورت زیر به دست می آید:

$$d^2(\tilde{A}, \tilde{B}) = (a - b)^2 + \frac{1}{6}(s_a - s_b)^2$$

۲. برآورد ريج برای مدل رگرسیون فازی با متغیرهای مستقل و پاسخ فازی

در این مقاله می خواهیم مدل رگرسیونی خطی چند گانه زیر را به داده ها برازش دهیم:

$$\tilde{Y} = \beta_0 \oplus \beta_1 \tilde{X}_1 \oplus \dots \oplus \beta_k \tilde{X}_k$$

که در آن $\tilde{Y} = (\tilde{Y}_1, \tilde{Y}_2, \dots, \tilde{Y}_k)'$ بردار متغیرهای پاسخ و $\tilde{Y}_i = (b_i, \sigma_i)_T$ و $\tilde{X}_i = (X_{i1}, \dots, X_{in})$ و $\tilde{X}_{ij} = (a_{ij}, s_{ij})_T$ که $i = 1, \dots, k$ و $j = 1, \dots, n$ اعداد فازی مثلثی هستند و β_i اعداد حقیقی اند.

با استفاده از روش کمترین مربعات خطا پارامترهای مدل خطی فازی را برآورد می کنیم. در این روش به دنبال مینیم کردن تابع $M(\beta)$ نسبت به β هستیم:

$$M(\beta) = \sum_{i=1}^k d^2(\beta_0 \oplus \beta_1 \tilde{x}_{i1} \oplus \dots \oplus \beta_k \tilde{x}_{in}, \tilde{y}_i)$$

نکته ۱.۲. فرض می شود که مقادیر $\beta_i, i = 1, \dots, k$ مثبت اند.

باتوجه به تعریف ۴.۱ تابع $M(\beta)$ به صورت زیر بازنویسی می شود:

$$M(\beta) = \sum_{i=1}^k (\beta_0 + \beta_1 a_{i1} + \dots + \beta_k a_{in} - b_i)^2 + \frac{1}{6} \sum_{i=1}^k (\beta_1 s_{i1} + \dots + \beta_k s_{in} - \sigma_i)^2$$

در ادامه با مشتق گیری از تابع $M(\beta)$ نسبت به β_i داریم:

$$\begin{aligned} k\beta_0 + \beta_1 \sum_{i=1}^k (a_{i1} a_{in} + \frac{1}{6} s_{i1} s_{in}) \\ + \dots + \beta_k \sum_{i=1}^k (a_{i1} a_{in} + \frac{1}{6} s_{i1} s_{in}) \\ = \sum_{i=1}^k (b_i a_{in} + \frac{1}{6} \sigma_i s_{in}) \end{aligned}$$

و حال با حل معادلات فوق به معادله ماتریسی زیر می رسم:

$$[X'_a X_a + \frac{1}{6} X'_s X_s] \beta = X'_a b + X'_s \sigma \quad (۱.۲)$$

که ماتریسهای X_s و X_a به ترتیب از بردارهای مراکز و پهنای متغیرهای مستقل تشکیل شده اند و بردارهای b و σ به ترتیب بردارهای مراکز و پهنای متغیرهای پاسخ به صورت زیر هستند:

$$b = (a_1, a_2, \dots, a_n)', \sigma = (\sigma_1, \sigma_2, \dots, \sigma_n)'$$

$$X_a = \begin{pmatrix} 1 & a_{11} & a_{21} & \dots & a_{k1} \\ 1 & a_{12} & a_{22} & \dots & a_{k2} \\ \dots & \dots & \dots & \dots & \dots \\ 1 & a_{1n} & a_{2n} & \dots & a_{kn} \end{pmatrix}$$

$$X_s = \begin{pmatrix} 0 & s_{11} & s_{21} & \dots & s_{k1} \\ 0 & s_{12} & s_{22} & \dots & s_{k2} \\ \dots & \dots & \dots & \dots & \dots \\ 0 & s_{1n} & s_{2n} & \dots & s_{kn} \end{pmatrix}$$

بنابراین برآورد β به صورت زیر است:

$$\hat{\beta} = [X'_a X_a + \frac{1}{6} X'_s X_s]^{-1} [X'_a b + X'_s \sigma]$$

اگر در معادله ۱.۲ ماتریس $X'_a X_a + \frac{1}{6} X'_s X_s$ رتبه کامل نباشد و یا کوچکترین مقدار ویژه ی آن بسیار نزدیک به صفر باشد، وارون پذیر نخواهد بود و برآورد کمترین مربعات خطا به دست

نخواهد آمد. در این مقاله از روش رگرسیون ریج برای براورد پارامترها هنگام بروز این مشکل استفاده شده است. برای براورد پارامترهای مدل خطی چندگانه فازی به روش ریج، تغییراتی در تابع $M(\beta)$ به صورت زیر ایجاد می شود: [۵، ۲، ۱]

$$M(\beta) = \gamma \sum_{i=1}^k \beta_i^2 + \sum_{i=1}^k (\beta_0 + \beta_1 a_{i1} + \dots + \beta_k a_{in} - b_i)^2 + \frac{1}{6} \sum_{i=1}^k (\beta_0 + \beta_1 s_{i1} + \dots + \beta_k s_{in} - \sigma_i)^2$$

که مقدار γ پارامترایی گفته می شود. با مشتق گیری از معادله فوق نسبت به β_i داریم:

$$\begin{aligned} \gamma \beta_k + k \beta_0 \left(1 + \frac{1}{6}\right) + \beta_1 \sum_{i=1}^k (a_{i1} a_{in} + \frac{1}{6} s_{i1} s_{in}) \\ + \dots + \beta_k \sum_{i=1}^k (a_{i1} a_{in} + \frac{1}{6} s_{i1} s_{in}) \\ = \sum_{i=1}^k (b_i a_{in} + \frac{1}{6} \sigma_i s_{in}) \end{aligned}$$

برآورد ریج برای پارامترهای β به صورت زیر است:

$$\hat{\beta}_R = [X'_a X_a + \frac{1}{6} X'_s X_s + \gamma I]^{-1} [X'_a b + X'_s \sigma] \quad (2.2)$$

مقادیر ماتریسهای X_a و X_s و بردارهای b و σ همان مقادیر معرفی شده در برآورد کمترین مربعات است. با توجه به نکته ۱.۲ اگر مقدار برآورد برای β_j منفی به دست آمد در معادله رگرسیونی، از ابتدا آن را منفی در نظر گرفته و دوباره روند برآوردیابی را تکرار می کنیم.

مثال ۲.۲. برای داده های جدول ۱ معادله رگرسیونی به شکل زیر است:

$$\tilde{Y} = \beta_0 \oplus \beta_1 \tilde{X}_1 \oplus \beta_2 \tilde{X}_2 \quad (3.2)$$

ماتریسهای X_a, X_s به شکل زیر هستند.

$$X_a = \begin{pmatrix} 1 & 2.3 & 3.6 \\ 1 & 3.2 & 5.4 \\ 1 & 3.5 & 6.0 \\ \vdots & \vdots & \vdots \\ 1 & 9.6 & 18.0 \end{pmatrix}$$

جدول ۱: مقادیر متغیرهای مستقل و پاسخ

	$\tilde{Y}$	$\tilde{X}_2$	$\tilde{X}_1$
۱	(۱۲/۵،۵/۰)	(۳/۶،۰/۱۸)	(۲/۳،۰/۲۳)
۲	(۱۷/۵،۵/۰)	(۵/۴،۰/۲۷)	(۳/۲،۰/۳۲)
۳	(۱۸/۵،۵/۰)	(۶/۰،۰/۳۰)	(۳/۵،۰/۳۵)
۴	(۲۲/۰،۵/۰)	(۷/۴،۰/۳۷)	(۴/۲،۰/۴۲)
۵	(۲۴/۰،۵/۰)	(۸/۲،۰/۴۱)	(۴/۶،۰/۴۶)
۶	(۲۶/۵،۵/۰)	(۹/۲،۰/۴۶)	(۵/۱،۰/۵۱)
۷	(۲۸/۰،۵/۰)	(۹/۸،۰/۴۹)	(۵/۴،۰/۵۴)
۸	(۳۱/۵،۵/۰)	(۱۱/۲،۰/۵۶)	(۶/۱،۰/۶۱)
۹	(۳۳/۵،۵/۰)	(۱۲/۰،۰/۶۰)	(۶/۵،۰/۶۵)
۱۰	(۳۷/۰،۵/۰)	(۱۳/۴،۰/۶۷)	(۷/۲،۰/۷۲)
۱۱	(۳۸/۵،۶/۰)	(۱۴/۰،۰/۷۰)	(۷/۵،۰/۷۵)
۱۲	(۴۱/۵،۶/۰)	(۱۵/۲،۰/۷۶)	(۸/۱،۰/۸۱)
۱۳	(۴۳/۵،۶/۰)	(۱۶/۰،۰/۸۰)	(۸/۵،۰/۸۵)
۱۴	(۴۶/۵،۶/۰)	(۱۷/۲،۰/۸۶)	(۹/۱،۰/۹۱)
۱۵	(۴۹/۰،۶/۰)	(۱۸/۰،۰/۹۰)	(۹/۶،۰/۹۶)

$$X_s = \begin{pmatrix} 0 & 0.23 & 0.18 \\ 0 & 0.32 & 0.27 \\ 0 & 0.35 & 0.30 \\ \vdots & \vdots & \vdots \\ 0 & 0.96 & 0.90 \end{pmatrix}$$

با تشکیل ماتریس $X'_a X_a + \frac{1}{6} X'_s X_s$ ، کوچکترین مقدار ویژه آن نزدیک به صفر است. برای برآورد پارامترها از روش ریج استفاده می‌کنیم. برای تعیین مقدار γ با توجه به شیوه‌ی معرفی شده در [۳] عمل می‌کنیم. مقداری انتخاب می‌شود که باعث ثبات در پارامترهای برآورد شده گردد. با توجه به جدول ۲ مقدار $\gamma = 0.002$ در نظر گرفته ایم. با توجه به معادله ۲.۲ برآورد ریج به دست آمده به صورت زیر است:

$$\hat{B}_R = \begin{pmatrix} -36/63 \\ 83/76 \\ -39/58 \end{pmatrix}$$

چون مقدار $\hat{\beta}_0, \hat{\beta}_2$ منفی به دست آمده است، ضریب آنها را در معادله رگرسیونی ۳.۲ منفی در نظر گرفته و برآورد یابی را دوباره انجام می‌دهیم. در نتیجه برآورد پارامترها به صورت زیر است:

$$\hat{B}_R = \begin{pmatrix} 35/90 \\ 5/57 \\ 3/01 \end{pmatrix}$$

جدول ۲: انتخاب مقدار γ

γ	$\hat{\beta}_0$	$\hat{\beta}_1$	$\hat{\beta}_2$
۰/۲	۳۵/۹۱	۵/۶۱	۳/۰۳
۰/۰۲	۳۵/۹۰	۵/۵۸	۳/۰۱
۰/۰۰۰۱	۳۵/۹۰	۵/۵۸	۳/۰۱
۰/۰۰۰۰۲	۳۵/۹۰	۵/۵۷	۳/۰۱

در ادامه برآورد کمترین مربعات پارامترهای مدل ۳.۲ محاسبه شده است:

$$\hat{B} = \begin{pmatrix} 35/97 \\ 5/60 \\ 2/00 \end{pmatrix}$$

براساس معیار نیکویی برازش معرفی شده در [۴] و با توجه به جدول زیر مقادیر نیکویی برازش همگی به مقدار یک نزدیکند و در نتیجه رگرسیون ریج به خوبی عمل کرده است.

جدول ۳: رگرسیون ریج و معیار نیکویی برازش

مقادیر برازش داده شده براساس رگرسیون ریج	معیار نیکویی برازش
(۲۳/۵۲، ۱۸/۹۱)	۰/۸۰
(۲۵/۹۴، ۱۹/۰۳)	۰/۸۸
(۲۶/۷۵، ۱۹/۰۷)	۰/۸۸
(۲۸/۶۳، ۱۹/۱۷)	۰/۹۲
(۲۹/۷۱، ۱۹/۲۲)	۰/۹۴
(۳۱/۰۵، ۱۹/۲۹)	۰/۹۶
(۳۱/۸۶، ۱۹/۳۳)	۰/۹۷
(۳۳/۷۴، ۱۹/۴۳)	۰/۹۹
(۳۴/۸۲، ۱۹/۴۸)	۰/۹۹
(۳۶/۷۰، ۱۹/۵۸)	۰/۹۹
(۳۷/۵۱، ۱۹/۶۲)	۰/۹۹
(۳۹/۱۲، ۱۹/۷۰)	۰/۹۹
(۴۰/۲۰، ۱۹/۷۵)	۰/۹۸
(۴۱/۸۱، ۱۹/۸۴)	۰/۹۶
(۴۲/۸۹، ۱۹/۸۹)	۰/۹۷

مراجع

1. Marvin, H.J.(1998).Improving Efficiency By Shrinkage, The James-Stein and Ridge regression Estimators.

2. Saunders, C., Gammerman, A. (1998). Ridge regression learning algorithm in dual variable In. proceedings of the 15th international conference on Machine Learning, 515-521.
3. Vapnik, V.N. (1998). Statistical Learning Theory, Springer, New York.
4. Xu, R., Li, C. (2001). Multidimensional Least-squares fitting with a Fuzzy model, Suzzy Sets and Systems 119, 215-223.
5. Hong, D.H. and Hwang, C. (2004). Ridge regression procedures for fuzzy models using triangular fuzzy numbers. International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems, 12: 145-159.



انتخاب مدل مناسب در رگرسیون لوجستیک فازی

فاطمه سلمانی^{۱*}، سید محمود طاهری^۲، علیرضا ابدی^۳، و حمید علوی مجد^۴

^۱ عضو هیئت علمی، دانشکده بهداشت، دانشگاه علوم پزشکی بیرجند
salmany_fatemeh@yahoo.com

^۲ عضو هیئت علمی، دانشکده فنی، دانشگاه تهران
sm_taheri@ut.ac.ir

^۳ عضو هیئت علمی، گروه پزشکی اجتماعی، دانشکده پزشکی، دانشگاه علوم پزشکی شهید بهشتی، تهران
alirezaabadi@gmail.com

^۴ عضو هیئت علمی، گروه آمار زیستی، دانشکده پیراپزشکی، دانشگاه علوم پزشکی شهید بهشتی، تهران
alavimajd@gmail.com

چکیده. وقتی تعداد متغیرهای شرکت کننده در مدل رگرسیون زیاد باشد، انتخاب بهترین مدل یکی از چالش‌ها است و همواره این سوال مطرح است که کدام زیرمجموعه از پیشگوها می‌توانند به بهترین صورت الگوی پاسخ را پیش‌بینی کنند و با چه فرآیندی می‌توان به این زیرمجموعه بهینه دست پیدا کرد. در این مقاله سعی شده است که به این سوالات در حالت رگرسیون لوجستیک فازی با پیشگوی فازی، پاسخ داده شود.

۱. انگیزه و اهمیت

در بسیاری از مطالعات، مجموعه متغیرهایی که می‌بایست در مدل رگرسیونی گنجانده شوند از پیش تعیین شده نیستند، و غالباً بخش اولیه تجزیه و تحلیل شامل انتخاب این متغیرها می‌باشد. در پاره‌ای موارد، انتخاب متغیرهای مستقل برای ورود به مدل بر مبنای ملاحظات بالینی و یا اصول خاصی است، در چنین مواقعی مسأله گزینش متغیر مطرح نیست. اما، در مواقعی که نظریه واضحی موجود نیست، گزینش متغیرها برای مدل رگرسیون موضوع مهمی خواهد بود. معمولاً در انتخاب مدل، دو معیار متقابل دخالت دارند.

۱- بدین منظور که معادله برای هدف‌های پیشگویی مفید باشد، ممکن است مایل باشیم مدل شامل حداکثر متغیرهای پیشگو باشد، به قسمی که بتوان مقادیر برازنده شده معتبر را تعیین کرد.
۲- به علت هزینه دستیابی به اطلاعات درباره تعداد زیاد متغیر پیشگو و ارائه آن‌ها، ممکن است مایل باشیم معادله شامل حداقل تعداد متغیر پیشگو باشد.

مصالحه بین این دو معیار، چیزی است که معمولاً آن را انتخاب بهترین معادله رگرسیون می‌نامند. شیوه یکتایی برای این کار وجود ندارد. حتی معیارهای مختلف انتخاب در یک مسأله خاص، لزوماً به یک مدل رگرسیونی یکتا و منحصر به فردی منجر نمی‌شوند [۱].

واژگان کلیدی. رگرسیون لوجستیک فازی، شاخص نیکویی برازش، انتخاب مدل.
* سخنران

۲. رگرسیون لوجستیک فازی

موضوع اصلی این مقاله، در حوزه آمار زیستی است. لذا بحث را با مساله‌ای در این حوزه طرح می‌کنیم. شرایطی را در نظر بگیرید که وضعیت ابتلا به بیماری در افراد، مبهم است. به عبارت دیگر در تشخیص بیماری یا حالت مورد نظر، ابهام وجود دارد. در مدل‌های آماری این نمونه‌های نادقیق از نمونه اصلی حذف می‌شوند. در مواردی که با چنین نمونه‌هایی سروکار داریم، از روش‌های معمول آماری نمی‌توان برای مدلبندی استفاده کرد و اساساً به دلیل نادقیق بودن متغیر وابسته، نمی‌توان احتمال $P(y_i = 1) = \pi_i$ را تعریف کرد. مدلی که در چنین موقعیتی پیشنهاد می‌شود، یک مدل فازی است که در آن به جای احتمال رویداد مورد نظر، امکان وقوع آن رویداد بر اساس عوامل خطر مدل‌بندی می‌شود. در صورتی که درجه سازگاری بیمار با شرایط بیماری به صورت واژه‌های زبانی مانند خیلی زیاد، زیاد، متوسط، کم و خیلی کم بیان شود، می‌توان به ازاء هر واژه زبانی یک مجموعه فازی با تابع عضویتی بر بازه (۰ و ۱) توسط فرد خبره (پزشک) تعریف کرد. اگر امکان وجود رویداد مورد نظر برای فرد i ام به صورت یک واژه زبانی بیان و با $\tilde{\mu}_i$ نشان داده شود، می‌توان هر واژه زبانی را به صورت یک مجموعه فازی با تابع عضویتی بر بازه (۰ و ۱) تعریف کرد. آنگاه نسبت $\frac{\tilde{\mu}_i}{1-\tilde{\mu}_i}$ ، بخت امکانی خصیصه موردنظر در فرد i ام است. این نسبت در واقع، امکان داشتن خصیصه مورد نظر به عدم داشتن آن را برای هر فرد نشان می‌دهد. با این توضیح برای تعریف مدل رگرسیون لوجستیک فازی، از تبدیل لگاریتم بخت امکانی $\tilde{W}_i = \ln(\frac{\tilde{\mu}_i}{1-\tilde{\mu}_i})$ ، $i = 1, 2, \dots, n$ به عنوان متغیر پاسخ استفاده می‌شود. به طوری که تبدیل بر روی یک مجموعه فازی صورت می‌گیرد. در این حالت با استفاده از اصل گسترش در مجموعه‌های فازی، تابع عضویت بخت امکانی محاسبه می‌شود [۴].

برای مدل‌بندی بین یک پاسخ فازی دوحالته مبنا و مجموعه‌ای از پیشگوهای فازی، بردار مشاهدات به صورت $(\tilde{x}_{i0}, \tilde{x}_{i1}, \dots, \tilde{x}_{ip-1}, \tilde{\mu}_i)$ در نظر گرفته می‌شود. به طوری که $\tilde{\mu}_i$ امکان وقوع رویداد مورد نظر برای فرد i ام که در ارتباط با یکی از عبارات زبانی متغیر پاسخ هستند، و $\tilde{x}_{ij}$ مقدار متغیر فازی j ام برای فرد i ام است.

$$\tilde{W}_i = \ln \frac{\tilde{\mu}_i}{1 - \tilde{\mu}_i} = b_0 + b_1 \tilde{x}_{i1} + \dots + b_{i(p-1)} \tilde{x}_{i(p-1)}, i = 1, 2, \dots, n$$

در این مدل $b_j, j = 0, 1, 2, \dots, p-1$ پارامترهای دقیق مدل است [۵].

۱.۲. برآورد نقطه‌ای پارامترهای مدل لوجستیک فازی با پیشگوی فازی. برای برآورد پارامترهای مدل با استفاده از کمترین مجموع مربعات، فاصله بین مقادیر مشاهده شده $\tilde{W}_i$ و مقادیر مشاهده شده $\tilde{w}_i$ باید کمینه شود. برای محاسبه فاصله بین دو عدد فازی از تعریف زیر استفاده می‌شود. فرض می‌کنیم که $\tilde{A}$ و $\tilde{B}$ دو عدد فازی مثلثی و $(\tilde{A})_\alpha = [a_1(\alpha), a_2(\alpha)]$ و $(\tilde{B})_\alpha = [b_1(\alpha), b_2(\alpha)]$ به ترتیب α -برش‌های $\tilde{A}$ و $\tilde{B}$ باشند. فاصله بین دو عدد فازی $\tilde{A}$ و $\tilde{B}$ بر پایه تابع وزنی $f(\alpha)$ به صورت زیر تعریف می‌شود:

$$d(\tilde{A}, \tilde{B}) = \left[\int_0^1 f(\alpha) d^2((\tilde{A})_\alpha, (\tilde{B})_\alpha) d\alpha \right]^{\frac{1}{2}}$$

و

$$d^2 \left(\left(\tilde{A} \right)_\alpha, \left(\tilde{B} \right)_\alpha \right) = [a_1(\alpha) - b_1(\alpha)]^2 + [a_2(\alpha) - b_2(\alpha)]^2$$

که در آن، $f(\alpha)$ یک تابع صعودی بر روی بازه $[0, 1]$ می باشد، به طوری که برای آن $f(0) = 0$ مقدار $\int_0^1 f(\alpha) d\alpha = \frac{1}{2}$. مقدار $d \left(\left(\tilde{A} \right)_\alpha, \left(\tilde{B} \right)_\alpha \right)$ فاصله بین α -برش های اعداد فازی $\tilde{A}$ و $\tilde{B}$ را اندازه می گیرد [۳].
در مدل لوجستیک فازی با پیشگوی فازی به فرم

$$\tilde{W}_i = \log \frac{\tilde{\mu}_i}{1 - \tilde{\mu}_i} = b_0 + b_1 \tilde{x}_{i1} + \dots + b_{p-1} \tilde{x}_{i,p-1}$$

اگر $\tilde{\mu}_i = (m_i, d_i^L, d_i^R)$ مقادیر مشاهده شده و $b_j, j = 0, 1, \dots, (p-1)$ ضرایب دقیق مدل $\tilde{X} = (\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_{(p-1)})$ و مقدار k و Z به صورت

$$k = \int_0^1 \alpha(\alpha-1) \ln \frac{(\alpha-1)d_i^L + m_i}{1 - [(\alpha-1)d_i^L + m_i]} d\alpha + \int_0^1 \alpha(1-\alpha) \ln \frac{(1-\alpha)d_i^R + m_i}{1 - [(1-\alpha)d_i^R + m_i]} d\alpha,$$

$$z = \int_0^1 \alpha \ln \frac{(\alpha-1)d_i^L + m_i}{1 - [(\alpha-1)d_i^L + m_i]} d\alpha + \int_0^1 \alpha \ln \frac{(1-\alpha)d_i^R + m_i}{1 - [(1-\alpha)d_i^R + m_i]} d\alpha,$$

باشند، آنگاه برآورد کمترین مربعات خطا مدل برابر است با

$$\underline{b} = \left[(X'X) + \frac{s's}{6} \right]^{-1} [X'z + s'k]$$

[۵].

۳. معیارهای انتخاب مدل

در مدل رگرسیون لوجستیک فازی، با در نظر گرفتن فاصله ای که بیان شده است، مجموع مربعات خطا به صورت زیر تعریف می شود

$$SSE^F = \sum_{i=1}^m \left(d(\tilde{w}_i, \tilde{W}_i) \right)^2 = \sum_{i=1}^n \int_0^1 \left[\left(\frac{\ln \frac{(\alpha-1)d_i^L + m_i}{1 - (\alpha-1)d_i^L - m_i}}{(\alpha-1)b_i s_{i1} - b_0 - b_1 x_{i1}} \right)^2 + \left(\frac{\ln \frac{(1-\alpha)d_i^R + m_i}{1 - (\alpha-1)d_i^R - m_i}}{(1-\alpha)b_i s_{i1} - b_0 - b_1 x_{i1}} \right)^2 \right] dx$$

میانگین مجموع مربعات خطا را با استفاده از تقسیم مجموع مربعات خطا بر $(n-p-1)$ محاسبه می شود.

$$MSE^F = \frac{SSE^F}{n-p-1}$$

AIC^F ، این معیار با تعمیم حالت آکائیکه غیرفازی به فازی به صورت زیر معرفی میشود:

$$AIC^F = n \ln(SSE^F) - n \ln(n) + 2(p+1)$$

که در آن SSE^F مجموع مربعات خطا، n حجم نمونه و p تعداد پارامترهای مدل است. C_P^F (شاخص مالوز در محیط فازی) کمیت دیگری است که در اینجا معرفی میشود که با تعمیم حالت غیرفازی به فازی، به صورت زیر تعریف می‌شود:

$$C_P^F = \frac{SSE_{Model}^F}{SSE_{saturated.Model}^F} + 2(p+3) - n$$

که در آن SSE_{Model}^F مجموع مربعات خطای مدل مورد بررسی است، $SSE_{saturated.Model}^F$ مجموع مربعات مدل شامل همه متغیرهای پیشگوست، n حجم نمونه و p تعداد پارامترهای مدل تحت بررسی است.

۴. روش‌های انتخاب مدل

۱.۴. تمام رگرسیون‌های ممکن. روش تمام رگرسیون‌های ممکن، شیوه‌ای نسبتاً خسته‌کننده و بدون داشتن یک ابزار سریع محاسبه، انجام آن غیرممکن است. در این شیوه ابتدا لازم است هر معادله رگرسیونی ممکن که شامل عرض از مبدا و هر تعداد از p متغیر پیشگوست، برازش داده شود؛ بنابراین $2^p - 1$ حالت وجود دارد؛ هر معادله رگرسیونی بر اساس معیاری انتخاب می‌شود. اما از آنجایی که ارزیابی کلیه رگرسیون‌های ممکن، به حجم زیاد محاسبه نیاز دارد، باید روش‌هایی ارائه شوند که صرفاً تعداد کمی از مدل‌های رگرسیون دارای زیرمجموعه‌ای از متغیرها را از طریق افزودن یا کاستن در یک زمان بررسی کند. این روش‌ها معمولاً به روش‌های گام به گام معروف هستند، که در اینجا برای مدل فازی مورد نظر روش انتخاب پیشرو معرفی می‌شود.

۲.۴. معرفی روش پیشرو برای رگرسیون لوجستیک فازی. روش معرفی شده در این بخش (مشابه حالت غیرفازی) روشی است که در هر مرحله متغیرهایی انتخاب می‌شوند که معیار لازم جهت ورود به معادله رگرسیونی را داشته باشند. بنابراین برای انتخاب متغیر در هر مرحله نیازمند تعریف معیار جدیدی هستیم. در ادامه تعریف این معیار ارائه داده میشود.

معیار Level of Efficacy (LE)

این معیار در جهت تصمیمگیری درباره حضور یا عدم حضور متغیر در مدل معرفی می‌شود. بدین ترتیب که متغیر وارد شده به مدل بتواند حداقل LE درصد، مدل قبلی را بهبود بخشد. اگر سطح تأثیر متغیر وارد شده به مدل پایین باشد، منجر به ورود متغیرهای بیشتری در مدل می‌شود. این موضوع باعث پیچیدگی بیشتر مدل شده و تحلیل مدل را دشوارتر می‌کند و پیش بینی با ابهام بیشتری صورت می‌گیرد. از سویی دیگر افزایش تعداد متغیرها در مطالعات بالینی، مستلزم صرف هزینه بیشتری در جهت جمع آوری اطلاعات مربوط به متغیرها ست. بنابراین انتخاب این معیار تا حد زیادی بستگی به شرایط مطالعه دارد. در ادامه به بیان الگوریتم پیشرو برای انتخاب مدل مناسب می‌پردازیم.

الگوریتم:

این روش با این فرض شروع می‌شود که هیچ متغیر مستقلاً در مدل حضور ندارد و فقط عرض از مبدأ وجود دارد. اولین متغیر پیشگوی فازی که برای ورود به معادله انتخاب می‌شود، متغیری

است که بیشترین تغییرات مدل را توصیف کرده باشد. در گام بعدی دلیل کافی برای حضور هر متغیر در مدل، به شرطی است که سطح تأثیر حداقل ۰/۱ را داشته باشد. در این شیوه گامهای اساسی عبارتند از:

- ۱- مدل‌های تک متغیره برازش داده می‌شود و بهترین مدل از بین آن‌ها انتخاب می‌شود.
- ۲- متغیر بعدی وارد مدل می‌شود و MSE^F محاسبه می‌شوند. اگر درصد بهبود مدل بیشتر LE درصد باشد، متغیر در مدل حضور پیدا می‌کند.
- ۳- این روند تا جایی که افزودن متغیر به مدل، باعث تغییری در شاخص نشود، ادامه پیدا می‌کند. در ادامه برای روشن شدن روند انتخاب مدل از مطالعه شبیه سازی استفاده می‌شود.

۵. مطالعه شبیه سازی

ابتدا مجموعه داده‌ای با سه متغیر پیشگو و پنج متغیر زاید تولید می‌شود. برای آن‌که شبیه سازی وابسته به اندازه پارامترها نباشد، انتخاب ضرایب اصلی مدل، به صورت تصادفی از توزیع یکنواخت $U(0/1, 3)$ صورت می‌گیرد. و به ازای هر انتخاب ۱۰۰۰ بار شبیه سازی تکرار و ضرایب متغیرهای زاید صفر در نظر گرفته می‌شود. میانگین مراکز متغیرهای اصلی از توزیع نرمال و پهناها از توزیع یکنواخت $U(0, 1)$ انتخاب شد. مراکز متغیرهای اضافی به ترتیب از توزیع‌های $x_4^c \sim N(10, 2)$ ، $x_5^c \sim N(4, 4)$ ، $x_6^c \sim N(5, 8)$ ، $x_7^c \sim N(2, 10)$ و $x_8^c \sim N(9, 1)$ با پهناهای تولید شده از $U(0, 1)$ انتخاب شد.

۱۰۵. نتایج مطالعه شبیه سازی. همه مدل‌های ممکن

از آنجا که هشت متغیر پیشگوی فازی تولید شد، 2^8 ترکیب مدل مختلف وجود دارد که باید برازش داده شوند. نمودارهای ۱ تا ۳ نتایج شاخص‌های نیکویی برازش حاصل از این مدل‌ها را نشان می‌دهد. نمودارها نشان می‌دهند که کمترین مقدار MSE^F مربوط به مدل کامل است. همه مدل‌های با MSE^F کوچک، شامل سه متغیر مؤثر هستند. اما تفاوت بین میانگین مربعات خطا مدل با سه متغیر و مدل کامل مقدار $\frac{1}{25}$ قابل اغماض است.

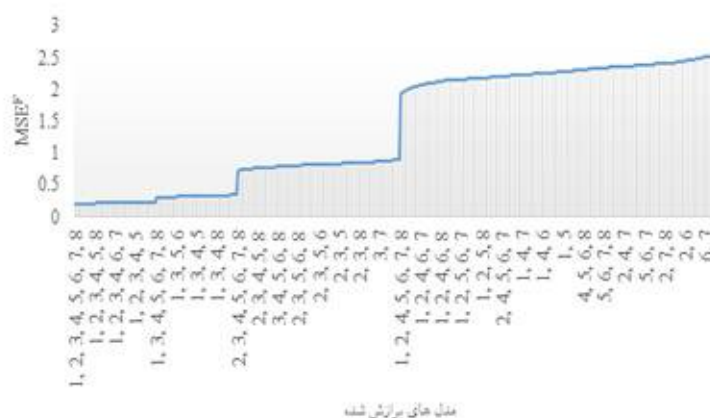
با استفاده از معیار آکائیکه بهترین مدل همان مدل با سه متغیر اصلی است و کمترین مقادیر شاخص مالوز مربوط به مدلی است که شامل سه متغیر اصلی هستند و فاصله شاخص‌های مدلی با تعداد متغیر بیشتر با مقدار اصلی بسیار ناچیز است.

- نتایج روش پیشرو

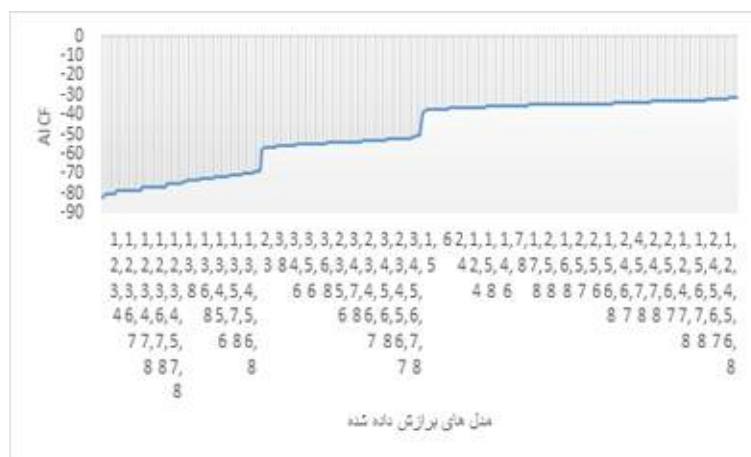
با در نظر گرفتن $LE=0/1$ نتایج در جدول ۱ ارائه شده است. بنابر نتایج جدول ۱ با روش حذف پیشرو، بهترین مدل شامل سه متغیر اصلی است که با نتایج شبیه سازی تطابق کامل دارد. در مرحله چهارم میزان بهبود مدل از حد مورد قبول کمتر است.

۶. مثال کاربردی

در این مقاله از اطلاعات مربوط به یک کارآزمایی بالینی استفاده شده است. این مطالعه در سال ۱۳۹۲ در بیمارستان آموزشی شفا دانشگاه علوم پزشکی کرمان انجام شده است. نمونه‌ها بیماران بستری در بخش ICU بیمارستان شفای شهر کرمان در سال ۱۳۹۱ می‌باشند که تحت عمل جراحی قلب باز قرار گرفته و دارای لوله قفسه سینه‌ای هستند. در این مطالعه ۱۲۸ بیمار مورد بررسی قرار گرفته‌اند. بیماران به طور تصادفی در دو گروه آزمون و شاهد قرار گرفتند. در هر دو گروه مراقبت‌های استاندارد بعد از عمل بر اساس پروتکل‌های قلبی و عروقی توسط جراحان



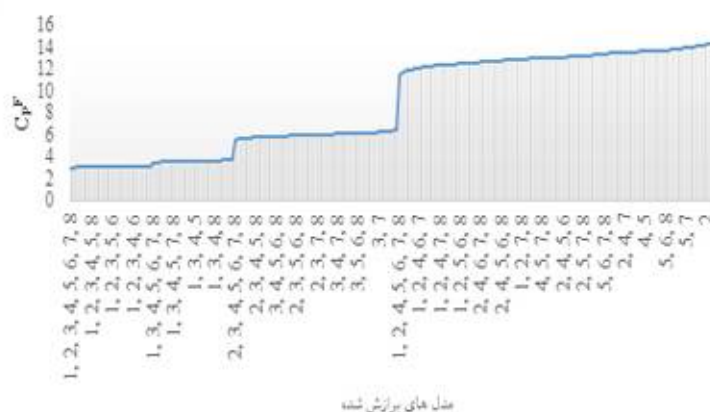
شکل ۱: تغییرات MSE^F در همه مدل‌های ممکن.



شکل ۲: تغییرات AIC^F در همه مدل‌های ممکن

و پرستاران انجام شده و سپس صدای فرد موردعلاقه وی که از روز قبل با ضبط صوت در اتاق انتظار ضبط شده با هدفون برای بیمار با شروع کار تا پایان آن گذاشته شده‌است. سپس شدت درد بیمار پنج دقیقه قبل، بلافاصله پس از کشیدن لوله قفسه سینه‌ای و پس از ۵ دقیقه سنجیده شده‌است.

۱۰۶. نتایج مطالعه کاربردی. بیماران شرکت‌کننده در این مطالعه میانگین سنی ۵۱/۱۵ با انحراف معیار ۱۳/۷ سال داشتند و ۶۸ درصد بیماران مرد بودند. در اینجا درد با استفاده از جملات ("بدون درد"، "اذیت‌کننده"، "کمی دردناک"، "دردناک"، "خیلی دردناک"، "غیرقابل تحمل") توصیف شده‌است. در همین راستا اعداد مثلی فازی به صورت زیر تعریف شدند.



شکل ۳: تغییرات C_p^F در همه مدل های ممکن

جدول ۱: مراحل انتخاب بهترین مدل

مرحله ۱	MSEF	مرحله ۲	بهبود مدل	مرحله ۳	بهبود مدل	مرحله ۴	بهبود مدل
$\tilde{x}_1$	۲/۳۹۳۷	$\tilde{x}_3$ و $\tilde{x}_1$	۵/۵۶۷۵	$\tilde{x}_1, \tilde{x}_2, \tilde{x}_3$	۱/۱۲۶۱	$\tilde{x}_1, \tilde{x}_2, \tilde{x}_3, \tilde{x}_4$	۳/۰۰۰۰
$\tilde{x}_2$	۲/۵۳۳۳	$\tilde{x}_3$ و $\tilde{x}_2$	۵/۰۵۰۶	$\tilde{x}_1, \tilde{x}_3, \tilde{x}_4$	۳/۰۰۰۳	$\tilde{x}_1, \tilde{x}_2, \tilde{x}_3, \tilde{x}_5$	۲۱/۰۰۰۰
$\tilde{x}_3$	۰/۹۱۷۸	$\tilde{x}_3$ و $\tilde{x}_4$	۴/۰۴۸۱	$\tilde{x}_1, \tilde{x}_3, \tilde{x}_5$	۴/۰۱۳۴	$\tilde{x}_1, \tilde{x}_2, \tilde{x}_3, \tilde{x}_6$	۳/۰۰۰۰
$\tilde{x}_4$	۲/۴۸۶۴	$\tilde{x}_3$ و $\tilde{x}_5$	۳/۰۳۸۸	$\tilde{x}_1, \tilde{x}_3, \tilde{x}_6$	۲/۰۱۲۲	$\tilde{x}_1, \tilde{x}_2, \tilde{x}_3, \tilde{x}_7$	۱۷/۰۰۰۰
$\tilde{x}_5$	۲/۵۱۳۹	$\tilde{x}_3$ و $\tilde{x}_6$	۲/۰۲۴۸	$\tilde{x}_1, \tilde{x}_3, \tilde{x}_7$	۴/۰۰۰۴	$\tilde{x}_1, \tilde{x}_2, \tilde{x}_3, \tilde{x}_8$	۵۴/۰۰۰۰
$\tilde{x}_6$	۲/۵۹۶۲	$\tilde{x}_3$ و $\tilde{x}_7$	۴/۰۴۳۷	$\tilde{x}_1, \tilde{x}_3, \tilde{x}_8$	۸/۰۰۰۸	-	-
$\tilde{x}_7$	۲/۵۶۸۵	$\tilde{x}_3$ و $\tilde{x}_8$	۶/۰۰۶۸	-	-	-	-
$\tilde{x}_8$	۲/۵۹۵۵	-	-	-	-	-	-

$$\mu_{None}(x) = \begin{cases} 1 - \frac{0.02-x}{0.01}, & 0.01 \leq x \leq 0.02, \\ 1 - \frac{x-0.02}{0.18}, & 0.02 < x \leq 0.2 \end{cases} \quad \mu_{Annoying}(x) = \begin{cases} 1 - \frac{0.2-x}{0.19}, & 0.01 \leq x \leq 0.2, \\ 1 - \frac{x-0.2}{0.2}, & 0.2 < x \leq 0.4. \end{cases}$$

$$\mu_{Uncomfortable}(x) = \begin{cases} 1 - \frac{0.4-x}{0.2}, & 0.2 \leq x \leq 0.4, \\ 1 - \frac{x-0.4}{0.2}, & 0.4 < x \leq 0.6. \end{cases} \quad \mu_{Dreadful}(x) = \begin{cases} 1 - \frac{0.6-x}{0.2}, & 0.4 \leq x \leq 0.6, \\ 1 - \frac{x-0.6}{0.2}, & 0.6 < x \leq 0.8. \end{cases}$$

$$\mu_{Horrible}(x) = \begin{cases} 1 - \frac{0.8-x}{0.2}, & 0.6 \leq x \leq 0.8, \\ 1 - \frac{x-0.8}{0.19}, & 0.8 < x \leq 0.99. \end{cases} \quad \mu_{Agonizing}(x) = \begin{cases} 1 - \frac{0.98-x}{0.18}, & 0.8 \leq x \leq 0.98, \\ 1 - \frac{x-0.98}{0.01}, & 0.98 < x \leq 0.99. \end{cases}$$

انتخاب مدل با استفاده از همه مدل های ممکن

مدل کامل، مدلی با همه پیشگوها به صورت زیر است.

$$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 sex_{ij} + b_3 age_{ij} + b_4 time_{ij} + b_5 location_{ij} + b_6 \tilde{\mu}_{ij-1}$$

در همه مدل‌ها متغیر پاسخ قبلی در مدل وارد می‌شود. در ادامه نتایج مربوط به برازش همه مدل‌ها آورده شده است.

جدول ۲: مرحله اول انتخاب بهترین مدل

شماره	مدل	MSE	C_p^F	AIC^F
۱	$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 \tilde{\mu}_{ij-1}$	۰/۰۰۳۷	-۲۳۶/۹۰۲	-۲۸۴۸/۱۳
۲	$\tilde{W}_{ij} = b_0 + b_1 sex_{ij} + b_2 \tilde{\mu}_{ij-1}$	۰/۰۰۴۴	-۲۳۶/۶۸۷	-۲۸۰۲/۴۲
۳	$\tilde{W}_{ij} = b_0 + b_1 age_{ij} + b_2 \tilde{\mu}_{ij-1}$	۰/۰۰۴۴	-۲۳۶/۶۸۷	-۲۸۰۰/۶۴
۴	$\tilde{W}_{ij} = b_0 + b_1 time_{ij} + b_2 \tilde{\mu}_{ij-1}$	۰/۰۰۳۶	-۲۳۶/۹۲	-۲۸۵۲/۳
۵	$\tilde{W}_{ij} = b_0 + b_1 location_{ij} + b_2 \tilde{\mu}_{ij-1}$	۰/۰۰۴۴	-۲۳۶/۶۷۸	-۲۸۰۰/۵۴

جدول ۳: مرحله دوم انتخاب بهترین مدل

مدل	MSE	C_p^F	AIC^F
$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 time_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۳۳	-۲۳۷/۰۱۲	-۲۸۷۳/۱۳
$\tilde{W}_{ij} = b_0 + b_1 age_{ij} + b_2 time_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۳۶	-۲۳۶/۹۱۶	-۲۸۴۹/۳۵
$\tilde{W}_{ij} = b_0 + b_1 sex_{ij} + b_2 time_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۳۶	-۲۳۶/۹۱۹	-۲۸۵۰/۰۹
$\tilde{W}_{ij} = b_0 + b_1 location_{ij} + b_2 time_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۳۶	-۲۳۶/۹۱۵	-۲۸۴۹/۲۷
$\tilde{W}_{ij} = b_0 + b_1 location_{ij} + b_2 group_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۳۷	-۲۳۶/۸۹۷	-۲۸۴۵/۰۳
$\tilde{W}_{ij} = b_0 + b_1 location_{ij} + b_2 age_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۴۵	-۲۳۶/۶۷۳	-۲۷۹۷/۷۲
$\tilde{W}_{ij} = b_0 + b_1 location_{ij} + b_2 sex_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۴۴	-۲۳۶/۶۸۵	-۲۷۹۹/۹۱
$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 age_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۳۷	-۲۳۶/۸۹۸	-۲۸۴۵/۲۴
$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 sex_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۳۷	-۲۳۶/۸۹۸	-۲۸۴۵/۱۸
$\tilde{W}_{ij} = b_0 + b_1 sex_{ij} + b_2 age_{ij} + b_3 \tilde{\mu}_{ij-1}$	۰/۰۰۴۴	-۶۸۲/۲۳۶	-۴۲/۲۷۹۹

با توجه به مقادیر همه شاخص‌ها، مدل‌هایی که شامل سه متغیر پیشگوی گروه، زمان و پاسخ پیشین، بهترین مدل‌ها هستند؛ و ورود متغیرهای پیشگوی اضافه باعث بهبود بیشتر مدل نمی‌شود. روش پیشرو

گام ۱. طبق جدول ۲، از بین مدل‌هایی که شامل یک متغیر پیشگو علاوه بر پاسخ پیشین هستند، بهترین مدل $\tilde{W}_{ij} = b_0 + b_1 time_{ij} + b_2 \tilde{\mu}_{ij-1}$ با میانگین کمترین مربعات ۰/۰۰۳۶ است. با در نظر گرفتن شاخص $LE=0/1$ ، مدلی در مرحله بعد مناسب است که حداقل ۰/۰۰۰۳۶ مدل را بهبود بخشد.

گام ۲. طبق جدول ۳، از بین مدل‌ها با دو متغیر پیشگو، تنها مدل شامل زمان، گروه و پاسخ پیشین ($\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 time_{ij} + b_3 \tilde{\mu}_{ij-1}$) با توجه به معیار $LE=0/0033$ MSE مورد پذیرش قرار می‌گیرد. در مرحله بعد حداقل بهبود در مدل باید ۰/۰۰۰۳۳ باشد.

جدول ۴: مرحله سوم انتخاب بهترین مدل

مدل	MSE	C_p^F	AIC^F
$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 age_{ij} + b_3 time_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۳۳	-۲۳۷/۰۰۸	-۲۸۷۰/۱۴
$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 sex_{ij} + b_3 time_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۳۳	-۲۳۷/۰۰۸	-۲۸۷۰/۱۶
$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 location_{ij} + b_3 time_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۳۳	-۲۳۷/۰۰۸	-۲۸۷۰/۱۲
$\tilde{W}_{ij} = b_0 + b_1 age_{ij} + b_2 location_{ij} + b_3 time_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۳۷	-۲۳۶/۹۱۱	-۲۸۴۶/۳۳
$\tilde{W}_{ij} = b_0 + b_1 age_{ij} + b_2 sex_{ij} + b_3 time_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۳۶	-۲۳۶/۹۱۵	-۲۸۴۷/۰۸
$\tilde{W}_{ij} = b_0 + b_1 sex_{ij} + b_2 location_{ij} + b_3 time_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۳۶	-۲۳۶/۹۱۵	-۲۸۴۷/۲
$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 location_{ij} + b_3 age_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۳۷	-۲۳۶/۸۹۳	-۲۸۴۲/۱۲
$\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 sex_{ij} + b_3 location_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۳۷	-۲۳۶/۸۹۳	-۲۸۴۲/۱۱
$\tilde{W}_{ij} = b_0 + b_1 group + b_2 age_{ij} + b_3 sex_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۳۷	-۲۳۶/۸۹۴	-۲۸۴۲/۲۵
$\tilde{W}_{ij} = b_0 + b_1 sex_{ij} + b_2 location_{ij} + b_3 age_{ij} + b_4 \tilde{\mu}_{ij-1}$	۰/۰۰۴۴	-۲۳۶/۶۷۹	-۲۷۹۶/۹

گام ۳. طبق جدول ۴، اضافه شدن متغیر جدید به مدل هیچ بهبودی در مدل ایجاد نمی‌کند. بنا بر هر دو روش بهترین مدل $\tilde{W}_{ij} = b_0 + b_1 group_{ij} + b_2 time_{ij} + b_3 \tilde{\mu}_{ij-1}$ است. جدول ۵ نتایج مربوط به مدل انتقال (مارکف) را نشان می‌دهد. روش برآورد پارامترها، کمترین مربعات فازی و روش انجام آزمون و فاصله اطمینان، بوت‌استرپ است.

جدول ۵: مرحله سوم انتخاب بهترین مدل

متغیر	ضریب	انحراف معیار	آماره آزمون	P-value	فاصله اطمینان بوت‌استرپ
عرض از مبدأ	۳/۸۶	۰/۲۶	۱۴/۷۶	<۰/۰۰۱	۴/۴۱
گروه	-۰/۷۸	۰/۰۹۷	-۸/۰۷	<۰/۰۰۱	-۰/۵۷
زمان	-۱/۷۴	۰/۱۲	-۱۵/۲۷	<۰/۰۰۱	-۱/۵۴
پاسخ پیشین	۱/۰۱	۰/۲۲	۴/۶۱	<۰/۰۰۱	۱/۴۵

مقدار آکائیکه مدل برابر ۲۸۷۳/۱۳- است.

۷. بحث

دی ارسو و سانتورو [۶] در مقاله‌ای با عنوان نیکویی برازش و انتخاب متغیر در رگرسیون خطی چندگانه فازی، به تشریح سه معیار ضریب تعیین، ضریب تعیین تعدیل شده و شاخص C_p پرداخته‌اند. نویسندگان مقاله از فاصله اقلیدسی در برآورد کمترین مربعات رگرسیون خطی فازی با پاسخ فازی و پیشگو دقیق، بهره جستند. آن‌ها با استفاده از مطالعه شبیه سازی صحت و دقت معیارهای پیشنهادی خود را بررسی نمودند. در این مطالعه حجم نمونه ۲۵ و ضرایب متغیرهای پیشگو ثابت بودند. مطالعه حاضر با در نظر گرفتن ضرایب مدل به صورت تصادفی، تعمیم پذیری

بیشتر نتایج را ایجاد کرده است. مطالعه شبیه سازی طراحی شده برای بررسی رویه انتخاب مدل، شامل هشت متغیر پیشگوی فازی بود که سه متغیر از بین این هشت متغیر، اصلی هستند. با توجه به نتایج MSE^F در رویکرد برازش همه مدل های ممکن، کوچکترین مقادیر MSE^F مربوط به مدل هایی است که همگی شامل سه متغیر تاثیرگذار هستند. اما فاصله بین MSE^F مدل اصلی تا کوچکترین MSE^F که شامل همه متغیرهای پیشگوست، تنها ۰/۰۲ است. معیار انتخاب مدل پیشنهاد شده در مطالعه حاضر، با تعمیم از حالت غیر فازی به فازی ارائه شده است. هر مدلی که کمترین مقدار معیار را اختیار کند، به عنوان مدل مورد نظر انتخاب میشود. در مطالعه حاضر، روش پیشرو به عنوان روشی جهت انتخاب مدل پیشنهاد شد؛ که در هر مرحله متغیرها تک تک به مدل اضافه شده و سپس با استفاده از مقداری که برای معیار سطح تاثیر در نظر گرفته می شود، مجوز ورود متغیر به مدل داده می شود. از میان مطالعاتی که به معرفی معیارهای نیکویی برازش و همچنین انتخاب متغیر مناسب جهت مدل بندی در زمینه رگرسیون فازی بوده، می توان به چند مطالعه اشاره کرد. پوراحمد و همکاران [۴] در مدل رگرسیون لوجستیک فازی با پیشگوی دقیق از معیار نیکویی برازش متوسط شاخص انطباق که توسط طاهری و کلکین نما [۷] معرفی شده بود، استفاده نمودند؛ اما رویکردی درباره روش های انتخاب مدل ارائه نداده اند. کامپوسو و فانیزی [۸] درباره روش های انتخاب مدل رگرسیون خطی فازی، روش پیش رونده را پیشنهاد کرده اند. در رویکرد مورد نظر، از شاخص ضریب تعیین برای ارزیابی مدل ها استفاده شده است. در هر مرحله از انتخاب مدل، مقدار ضریب تعیین با یک مقدار آستانه مقایسه می شود. یکی از معایبی که نویسندگان درباره روش پیشنهادی خود مطرح می کنند این است که متغیری که در یکی از مراحل انتخاب مدل، از الگوریتم خارج می شود، بار دیگر در ترکیب های بعدی مدل نمی تواند حضور پیدا کند. وجه تشابه مطالعه حاضر با این مطالعه در انتخاب روش پیشرو به عنوان روش انتخاب مدل مناسب است و همچنین استفاده از یک معیار مفروض جهت قضاوت در هر مرحله است با این تفاوت که معیار مورد نظر در مطالعه کامپوسو و فانیزی مانند مطالعه دی اورو، ضریب تعیین است در حالی که در مطالعه ما با معرفی معیار سطح تاثیر (LE) از میانگین مربعات خطا استفاده نموده اند. همچنین مزیت الگوریتم ارائه شده در مطالعه حاضر نسبت به مطالعه کامپوسو و فانیزی در این است که الگوریتم به گونه ای طراحی شده است که با توجه میزان بهبودی که ورود متغیر جدید در مدل ایجاد می کند، متغیر حذف شده در هر مرحله، امکان ورود مجدد به مدل را دارد. از آنجایی که برازش همه مدل های ممکن زمانی که تعداد متغیرهای پیشگو زیاد شوند، کار وقت گیری ست، روش پیشرو با تعداد برازش های کمتری به مدل نهایی دست می یابد. اگر سطح تاثیر متغیر وارد شده به مدل پایین باشد، منجر به ورود متغیرهای بیشتری در مدل می شود. این موضوع باعث پیچیدگی بیشتر مدل شده و تحلیل مدل را دشوارتر می کند و پیش بینی با ابهام بیشتری صورت می گیرد. از سویی دیگر افزایش تعداد متغیرها در مطالعات بالینی، مستلزم صرف هزینه بیشتری در جهت جمع آوری اطلاعات مربوط به متغیرها می شود.

مراجع

1. Rezaei, A. and A. Soltani (1998), *An introduction to applied regression analysis*, Isfahan, Isfahan university of technology.
2. Zadeh, L. A. (1975), *The concept of a linguistic variable and its application to approximate reasoning*, Information sciences 8(3):199-249.

3. Xu, R. and C. Li (2001), *Multidimensional least-squares fitting with a fuzzy model*, Fuzzy sets and systems .119(2): 215-223.
4. Pourahmad, S., S. M. T. Ayatollahi, S. M. Taheri and Z. H. Agahi (2011), *Fuzzy logistic regression based on the least squares approach with application in clinical studies*, Computers Mathematics with Applications 62(9):3353-3365.
5. Salmani, F., S. M. Taheri, J. H. Yoon, A. Abadi, H. A. Majd and A. Abbaszadeh (2011), *Logistic Regression for Fuzzy Covariates: Modeling, Inference, and Applications*, International Journal of Fuzzy Systems: 1-10.
6. D'Urso, P. and A. Santoro (2006), *Goodness of fit and variable selection in the fuzzy multiple linear regression*, Fuzzy Sets and Systems 157(19):2627-2647.
7. Taheri, S. M. and M. Kelkinnama (2008), *Fuzzy least absolute regression*, Intelligent Systems. 4th International IEEE Conference, IEEE.
8. Campobasso, F. and A. Fanizzi (2013), *Goodness of fit measures and model selection in a fuzzy least squares regression analysis*, Computational Intelligence. J. Kacprzyk. New York Springer. 7092:241-257.



شاخص‌های کارایی فرایند تعمیم یافته با اعداد فازی و کیفیت فازی

نگار شیخی^۱، محمد رضا ربیعی^۱، و عباس پرچمی^۲

^۱ شاهرود، دانشگاه صنعتی شاهرود، دانشکده علوم ریاضی، گروه آمار

sheikhi.negar1611@gmail.com

rabiei1354@yahoo.com

^۲ کرمان، دانشگاه شهید باهنر کرمان، دانشکده ریاضی و کامپیوتر، بخش آمار

parchami@uk.ac.ir

چکیده. شاخص کارایی فرایند آماری برای برآورد توانایی ذاتی یک فرایند تولیدی در تولید محصولات منطبق با مشخصات فنی محصول است. اولین نسل شاخص‌های کارایی فرایند (C_p) به صورت نسبت پهنای حدود مشخصه فنی به پهنای حدود تلورانس طبیعی معرفی گردید. C_{pm} به عنوان شاخص کارایی فرایند دیگر که شامل دو تغییر جزئی میانگین فرایند و انحراف میانگین فرایند از هدف می‌باشد نیز معرفی شد. این مقاله به تعمیم شاخص‌های C_p و C_{pm} در یک محیط فازی می‌پردازد. این شاخص‌های کارایی تعمیم یافته قادر به اندازه‌گیری کارایی فرایند فازی مبتنی بر کیفیت فازی هستند. فرمول محاسباتی سریع شاخص‌های کارایی فرایند تعمیم یافته زمانی که مشاهدات اعداد فازی مثلثی و نرمال نامتقارن هستند، محاسبه شده است که در آن کیفیت فازی بر اساس کران‌های فازی خطی و نمایی تعریف شده‌اند.

۱. مقدمه

شاخص کارایی فرایند^۱ در حقیقت نوعی خلاصه‌شده عددی است که میزان کارایی فرایند را نسبت به حدود مشخصه فنی^۲ ارزیابی می‌کند. چندین شاخص کارایی مانند C_p ، C_{pk} ، C_{pm} و C_{pmk} وجود دارند که برای تخمین کارایی فرایندهای نرمال (معمولاً با حجم بزرگ) استفاده می‌شوند. فرض کنید مقدار امید ریاضی و انحراف استاندارد متغیر تک بعدی X را به ترتیب با μ و σ نشان می‌دهیم. C_p منسوب به ژوران [۵] و C_{pm} منسوب به هسیانگ و تاگوچی [۴] به صورت زیر معرفی شده‌اند:

$$C_p = \frac{USL - LSL}{6\sigma} \quad (1.1)$$

و

$$C_{pm} = \frac{USL - LSL}{6\sqrt{E(X - T)^2}} \quad (2.1)$$

واژگان کلیدی. کیفیت فازی، اعداد فازی مثلثی و نرمال، شاخص کارایی فرایند، عملگرهای حسابی.
* سخنران

^۱Process Capability Index

^۲Specification Limit

به طوری که USL^3 حدود مشخصه بالایی و LSL^4 حدود مشخصه پایینی و T مقدار هدف است. در سال‌های اخیر، بعضی مقالات که تمرکز آن‌ها روی زمینه‌های مختلفی از شاخص‌های کارایی با استفاده از نظریه مجموعه فازی بودند منتشر شده است. این زمینه‌ها عبارت‌اند از: محاسبه شاخص‌های کارایی فازی زمانی که مشاهدات اعداد فازی‌اند [۷]، بررسی روی C_{pm} فازی و انتخاب بهترین گزینه [۳]، برآورد شاخص‌های کارایی با استفاده از روش برآورد فازی باکلی [۹]، تحلیل کارایی فرایند بر اساس اندازه‌های فازی و نمودارهای کنترل فازی [۶]، معرفی یک نسل جدید از شاخص‌های کارایی بر اساس اندازه‌های فازی [۱۰] و معرفی یک نسل جدید از شاخص‌های کارایی بر اساس حدود مشخصه فازی [۸]. در این مقاله، بر خلاف این مطالعات، شاخص‌های کارایی با فرض یک متر معنی‌دار تعمیم یافته‌اند به طوری که داده‌ها و حدود مشخصه‌ها فازی/مبهم در نظر گرفته شده‌اند. ساختار این مقاله به شرح زیر است:

پس از بحث در مورد مقدمات مجموعه فازی، بخش ۲ به تعریف حدود مشخصه فازی و محاسبه‌ی فاصله‌ی بین آن‌ها پرداخته است. در بخش ۳ چندین شاخص‌های کارایی تعمیم یافته بر اساس کیفیت فازی و داده غیرفازی مرور شده است. در بخش ۴ دو نسخه ساختار فوق‌العاده از C_p و C_{pm} معرفی می‌شود که قادر به ارائه اندازه‌های عددی برای میزان توانایی یک فرایند فازی-مقدار همراه با حدود مشخصه فازی است. یک مثال کاربردی برای نشان دادن کارایی شاخص‌های پیشنهادی در بخش ۵ ارائه شده است. بخش انتهایی بخش نتیجه‌گیری است.

۲. حدود مشخصات فازی

فرض کنید $\mathbb{R}$ مجموعه اعداد حقیقی و $F(\mathbb{R}) = \{\tilde{A} | \tilde{A} : \mathbb{R} \rightarrow [0, 1]\}$ مجموعه تمام زیرمجموعه‌های فازی از $\mathbb{R}$ باشد. α -برش $\tilde{A}$ یک مجموعه غیرفازی است که برای هر $\alpha \in (0, 1]$ با $\tilde{A}_\alpha = \{x \in \mathbb{R} | \tilde{A}(x) \geq \alpha\}$ نشان داده می‌شود.

تعریف ۱.۲. [۱] $\widetilde{USL}_L$ کران بالایی فازی $\widetilde{USL}_L \in F_U(\mathbb{R})$ و $\widetilde{LSL}_L$ کران پایینی فازی $\widetilde{LSL}_L \in F_L(\mathbb{R})$ را به ترتیب کران فازی بالایی و پایینی خطی می‌نامیم، هرگاه توابع عضویت آن‌ها به شکل زیر باشد:

$$\widetilde{USL}_L(x) = \begin{cases} 1 & x \leq u_1 \\ \frac{x-u_0}{u_1-u_0} & u_1 \leq x \leq u_0 \\ 0 & x > u_0 \end{cases} \quad (1.2)$$

و

$$\widetilde{LSL}_L(x) = \begin{cases} 0 & x \leq l_0 \\ \frac{x-l_0}{l_1-l_0} & l_0 \leq x \leq l_1 \\ 1 & x > l_1 \end{cases} \quad (2.2)$$

به طوری که $u_0, u_1, l_0, l_1 \in \mathbb{R}$ و $l_0 < l_1 \leq u_1 < u_0$. برای سهولت آن‌ها را به ترتیب با $\widetilde{USL}_L = (u_1, u_0)_{UL}$ و $\widetilde{LSL}_L = (l_0, l_1)_{LL}$ نشان می‌دهیم.

^۳Upper Specification Limit

^۴Lower Specification Limit

تعریف ۲.۲. کران بالایی فازی $\widetilde{USL}_E \in F_U(\mathbb{R})$ و کران پایینی فازی $\widetilde{LSL}_E \in F_L(\mathbb{R})$ را به ترتیب کران‌های فازی بالایی و پایینی نمایی می‌نامیم، هرگاه توابع عضویت آن‌ها به شکل زیر باشد:

$$\widetilde{USL}_E(x) = \begin{cases} 1 & x \leq u_1 \\ e^{-[(x-u_1)/s_u]^2} & x > u_1 \end{cases} \quad (۳.۲)$$

و

$$\widetilde{LSL}_E(x) = \begin{cases} 1 & x \geq l_1 \\ e^{-[(x-l_1)/s_l]^2} & x < l_1 \end{cases} \quad (۴.۲)$$

به طوری که $u_1, s_u, l_1, s_l \in \mathbb{R}$ و $l_1 \leq u_1$ و $s_u, s_l > 0$. که برای سهولت آن‌ها را به ترتیب $\widetilde{USL}_E = (u_1, s_u)_{UE}$ و $\widetilde{LSL}_E = (l_1, s_l)_{LE}$ نشان می‌دهیم.

اکنون، $F_U(\mathbb{R})$ و $F_L(\mathbb{R})$ به ترتیب مجموعه‌ای از همه کران‌های فازی بالایی و مجموعه‌ای از همه کران‌های فازی پایینی را مشخص می‌کنند.

با جایگزینی کران‌های فازی به جای غیر فازی می‌توان یک قضاوت موجه در تصمیم‌گیری روی فرایندهای تولید انتظار داشت. یکی دیگر از مزایای فرض کران‌های فازی به جای غیر فازی افزایش توانایی مهندسين برای تعريف USL و LSL است. از این رو مهندسين می‌توانند تولیداتی با اندازه‌های مشخصه پایین‌تر از LSL یا بالاتر از USL را با درجه عضویت بین $[0, 1]$ به عنوان انطباق تولیدات (تا رسیدن به استاندارد) بپذیرند. بنابراین معرفی و فرض کران‌های فازی می‌تواند تعداد زیادی از موارد غیر انطباق را به حداقل رساند و هم‌چنین هزینه‌های تولید را کاهش دهد.

تعریف ۳.۲. [۲] اگر $\widetilde{USL} \in F_U(\mathbb{R})$ و $\widetilde{LSL} \in F_L(\mathbb{R})$ ، آنگاه فاصله‌ی میان دو کران فازی $\widetilde{USL}$ و $\widetilde{LSL}$ را با نماد $D(\widetilde{USL}, \widetilde{LSL})$ نشان داده و به صورت زیر تعریف می‌کنیم

$$D(\widetilde{USL}, \widetilde{LSL}) = \int_0^1 g(\alpha)(u_\alpha - l_\alpha) d\alpha \quad (۵.۲)$$

که در آن

$$l_\alpha = \inf_x \{x \in \widetilde{LSL}_\alpha\} \quad \text{و} \quad u_\alpha = \sup_x \{x \in \widetilde{USL}_\alpha\}$$

از $\widetilde{USL}_\alpha = (-\infty, u_\alpha]$ و $\widetilde{LSL}_\alpha = [l_\alpha, +\infty)$ به دست آمده‌اند و g تابعی نانزولی روی بازه $[0, 1]$ می‌باشد به طوری که $g(0) = 0$ و $\int_0^1 g(\alpha) d\alpha = 1$ است (به عنوان مثال، به ازای $m = 0, 1, 2, \dots$ می‌توان g را به صورت $g(\alpha) = (m+1)\alpha^m$ در نظر گرفت).

توجه داشته باشید که در تعریف ۳.۲ $u_\alpha - l_\alpha$ تفاضل نقاط انتهایی راست و چپ از $-\alpha$ برش $\widetilde{USL}$ و $\widetilde{LSL}$ است. مقدار $g(\alpha)$ می‌تواند به عنوان وزنی از $u_\alpha - l_\alpha$ باشد و ویژگی نانزولی بودن g بدین معنی است که بالاترین مقدار α در تعیین تفاضل $\widetilde{USL}$ و $\widetilde{LSL}$ دارای اهمیت بیشتری است. مزیت عملگر D این است که می‌تواند مجموعه‌های $-\alpha$ برش مختلف را با وزن‌های متفاوت در محاسبه در برگیرد.

توجه ۴.۲. طبق نکته ۱ [۱]، در حالتی که حدود مشخصه به صورت غیر فازی با توابع مشخصه $I_{\{x; x \geq LSL\}}$ و $I_{\{x; x \leq USL\}}$ در نظر گرفته می‌شوند، می‌توان نشان داد که

$$D(\widetilde{USL}, \widetilde{LSL}) = \int_0^1 g(\alpha)(u_\alpha - l_\alpha) d\alpha = USL - LSL$$

قضیه ۵.۲. [۸] فرض کنید $\widetilde{USL}_L = (u_1, u_0)_{UL}$ و $\widetilde{LSL}_L = (l_0, l_1)_{LL}$ به ترتیب کران‌های فازی بالایی و پایینی خطی باشند. اگر $m = 0, 1, 2, \dots$ ، $g(\alpha) = (m+1)\alpha^m$ ، آن‌گاه

$$D(\widetilde{USL}_L, \widetilde{LSL}_L) = \frac{1}{m+2}[(m+1)(u_1 - l_1) + (u_0 - l_0)] \quad (۶.۲)$$

قضیه ۶.۲. [۸] فرض کنید $\widetilde{USL}_E = (u_1, s_u)_{UE}$ و $\widetilde{LSL}_E = (l_1, s_l)_{LE}$ به ترتیب کران‌های فازی بالایی و پایینی نمایی باشند. اگر $m = 0, 1, 2, \dots$ ، $g(\alpha) = (m+1)\alpha^m$ ، آن‌گاه

$$D(\widetilde{USL}_E, \widetilde{LSL}_E) = u_1 - l_1 + \sqrt{\pi/4(m+1)}(s_u + s_l) \quad (۷.۲)$$

۳. شاخص‌های کارایی فرایند بر اساس کیفیت فازی و داده فازی

با توجه به عدم ارتباط حدود مشخصه فنی و مقادیر مشاهده شده از مشخصه فنی در یک فرایند تولیدی، ابهام موجود در آن‌ها نیز به یک دیگر مرتبط نبوده و بنابراین در عمل، ممکن است با چهار حالت زیر مواجه شویم: (۱) کران‌ها غیر فازی و داده غیر فازی، (۲) کران‌ها فازی و داده غیر فازی، (۳) کران‌های غیر فازی و داده فازی (۴) کران‌های فازی و داده فازی. اکنون مرور مختصری بر عملگرها و تعاریف نظریه مجموعه فازی که در این مقاله به کار برده شده‌اند داریم.

تعریف ۱.۳. (الف) یک حالت خاص از اعداد فازی، اعداد فازی مثلثی نامتقارن هستند که به صورت زیر تعریف می‌شوند:

$$\tilde{T}(a, s^L, s^R)_{(x)} = \begin{cases} \frac{x-a+s^L}{s^L} & a - s^L \leq x \leq a \\ \frac{a+s^R-x}{s^R} & a < x \leq a + s^R \\ 0 & \text{در جاهای دیگر} \end{cases} \quad (۱.۳)$$

که به صورت نمادین با $T(x, s_x^L, s_x^R)$ نشان می‌دهیم.

(ب) یک حالت خاص دیگر از اعداد فازی، اعداد فازی نرمال نامتقارن هستند که به صورت زیر تعریف می‌شوند:

$$\tilde{N}(a, s_a^L, s_a^R)_{(x)} = \begin{cases} e^{-[(a-x)/s_a^L]^2} & x \leq a \\ e^{-[(x-a)/s_a^R]^2} & x > a \end{cases} \quad (۲.۳)$$

و به صورت نمادین با $N(a, s_a^L, s_a^R)$ نشان می‌دهیم.

عدد حقیقی a و اعداد حقیقی مثبت s^L و s^R به ترتیب مقدار نما و پهنای چپ و راست این دو نوع از اعداد فازی اند. همچنین $F_T(\mathbb{R})$ و $F_N(\mathbb{R})$ به ترتیب مجموعه‌ای از همه اعداد فازی مثلثی نامتقارن و اعداد فازی نرمال نامتقارن را مشخص می‌کنند. $T(a, 0, 0)$ و $N(a, 0, 0)$ تابع نشانگر a یعنی، $I_{\{a\}}$ هستند.

با الهام گرفتن از تعریف ۳.۲، به منظور اندازه‌گیری فاصله بین دو عدد فازی، یک عملگر حسابی روی اعداد فازی که بر اساس تابع وزنی $g(\alpha)$ است ارائه می‌دهیم. این فاصله به طور کامل در بخش ۴ برای معرفی نسل جدیدی از شاخص‌های کارایی بر اساس مشاهدات فازی و کیفیت فازی به کار برده خواهد شد.

تعریف ۲.۳. برای هر $\tilde{A}, \tilde{B} \in F(\mathbb{R})$

$$d(\tilde{A}, \tilde{B}) = \left(\int_0^1 \frac{g(\alpha)}{2} [\tilde{A}_\alpha - \tilde{B}_\alpha]^2 d\alpha \right)^{\frac{1}{2}} \quad (۳.۳)$$

فاصله بین $\tilde{A}$ و $\tilde{B}$ نامیده می‌شود، که در آن برای هر $\alpha \in [0, 1]$

$$\tilde{A}_\alpha - \tilde{B}_\alpha = \{[a_1(\alpha) - b_1(\alpha)]^2 + [a_2(\alpha) - b_2(\alpha)]^2\}^{\frac{1}{2}} \quad (۴.۳)$$

فاصله بین $\tilde{A}_\alpha = [a_1(\alpha), a_2(\alpha)]$ و $\tilde{B}_\alpha = [b_1(\alpha), b_2(\alpha)]$ اندازه‌گیری شده و شرایط تابع g در تعریف ۳.۲ آمده است.

تعریف ۳.۳. [۸] فرض کنید $\widetilde{USL}$ و $\widetilde{LSL}$ به ترتیب کران‌های فازی بالایی و پایینی باشند. برآورد شاخص‌های کارایی تعمیم یافته بر اساس مشاهدات فازی $\tilde{x}_i \in F(\mathbb{R})$ و $i = 1, \dots, n$ به صورت زیر است:

$$\hat{C}_{\tilde{p}} = \frac{D(\widetilde{USL}, \widetilde{LSL})}{6\sqrt{\frac{1}{n} \sum_{i=1}^n d^2(\tilde{x}_i, \tilde{x})}} \quad (۵.۳)$$

$$\hat{C}_{\tilde{pm}} = \frac{D(\widetilde{USL}, \widetilde{LSL})}{6\sqrt{d^2(\tilde{x}, I_{\{T\}}) + \frac{1}{n} \sum_{i=1}^n d^2(\tilde{x}_i, \tilde{x})}} \quad (۶.۳)$$

به طوری که $I_{\{T\}}$ تابع نشانگر مقدار هدف و $\tilde{x}$ میانگین مشاهدات فازی $\tilde{x}_i$ است.

توجه ۴.۳. برآورد شاخص‌های کارایی تعمیم یافته بر اساس مشاهدات فازی بسطی مناسب برای شاخص‌های کارایی معمولی است. از این رو با در نظر داشتن ملاحظه ۴.۲، وقتی که داده‌های مشاهده شده دقیق و حدود مشخصه فنی مربوطه با توابع نشانگر $I_{\{x:x \leq USL\}}$ و $I_{\{x:x \geq LSL\}}$ طراحی شوند، برآورد شاخص‌های تعمیم یافته $C_{\tilde{pm}}$ و $C_{\tilde{p}}$ منطبق بر برآورد شاخص‌های معمولی C_p و C_{pm} می‌شوند.

در دو قضیه زیر، قصد داریم تا چند فرمول محاسباتی سریع برای واریانس اعداد فازی مثلثی نامتقارن و نرمال نامتقارن به دست آوریم که در این بخش برای برآورد شاخص‌های کارایی بر اساس مشاهدات فازی استفاده خواهند شد.

قضیه ۵.۳. فرض کنید $\tilde{x}_i = T(x_i, s_{x_i}^L, s_{x_i}^R) \in F_T(\mathbb{R})$ و $i = 1, \dots, n$. بر اساس تابع وزنی $g(\alpha) = (m+1)\alpha^m$ و $m = 0, 1, 2, \dots$ ، واریانس اعداد فازی مثلثی نامتقارن $\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_n$ به صورت زیر است:

$$v_T(\tilde{x}) = \left[v(x) + \frac{1}{(m+2)(m+3)} [v(s_x^R) + v(s_x^L)] \right. \\ \left. + \frac{1}{m+2} [\text{cov}(x, s_x^R) - \text{cov}(x, s_x^L)] \right] \quad (۷.۳)$$

که در آن

$$v(s_x^L) = \frac{1}{n} \sum_{i=1}^n (s_{x_i}^L - \bar{s}_x^L)^2, v(x) = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$$

و

$$v(s_x^R) = \frac{1}{n} \sum_{i=1}^n (s_{x_i}^R - \bar{s}_x^R)^2$$

به ترتیب واریانس x_i ها، $s_{x_i}^L$ ها و $s_{x_i}^R$ ها هستند. و همچنین دو تساوی

$$\text{cov}(x, s_x^R) = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(s_{x_i}^R - \bar{s}_x^R)$$

و

$$\text{cov}(x, s_x^L) = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})(s_{x_i}^L - \bar{s}_x^L)$$

برقرارند.

برهان. با استفاده از تعریف تابع عضویت اعداد مثلثی نامتقارن (۱.۳) و تعریف ۲.۳ فاصله دو عدد فازی و انتخاب $g(\alpha) = (m+1)\alpha^m$ به ازای $m = 0, 1, 2, \dots$ به سادگی قابل محاسبه است. $\square$

قضیه ۶.۳. فرض کنید $\tilde{x}_i = N(x_i, s_{x_i}^L, s_{x_i}^R) \in F_N(\mathbb{R})$ و $i = 1, \dots, n$. بر اساس تابع وزنی $g(\alpha) = (m+1)\alpha^m$ و $m = 0, 1, 2, \dots$ ، واریانس اعداد فازی نرمال نامتقارن $\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_n$ به صورت زیر است:

$$v_N(\tilde{x}) = v(x) - \frac{1}{2(m+1)} [v(s_x^R) + v(s_x^L)] + \sqrt{\frac{\pi}{4(m+1)}} \\ \times [\text{cov}(x, s_x^R) - \text{cov}(x, s_x^L)] \quad (۸.۳)$$

و شرایط این قضیه مشابه قضیه ۵.۳ است.

برهان. با استفاده از تعریف تابع عضویت اعداد فازی نرمال نامتقارن (۲.۳) و تعریف ۲.۳ و انتخاب $g(\alpha) = (m+1)\alpha^m$ به ازای $m = 0, 1, 2, \dots$ به سادگی قابل محاسبه است. $\square$

نتیجه ۷.۳. [۲] کسرهای

$$\sqrt{\frac{\pi}{4(m+1)}} [\text{cov}(x, s_x^R) - \text{cov}(x, s_x^L)] \text{ و } \frac{1}{m+2} [\text{cov}(x, s_x^R) - \text{cov}(x, s_x^L)]$$

معمولاً عددی نسبتاً کوچک می‌باشند و به دلیل عدم تقارن موجود در اعداد مثلثی و نرمال در قضیه ۵.۳ و ۶.۳ پدیدار شده است.

مثال ۸.۳. مثال کاربردی مربوط به تولید پیستون در محیط صنعتی قونیه از مرجع [۶] را در نظر بگیرید. که در آن مشخصه کیفیت قطر داخلی رینگ پیستون می‌باشد. می‌خواهیم اعداد فازی مثلی با همان مقادیر اصلی اما با پهنای متفاوتی را به کار ببریم. هسته‌ی این اعداد برابر هسته‌ی اعداد کایا و قهرمان در نظر گرفته شده‌اند ولی پهنای این اعداد، یعنی $s_{x_i}^L$ و $s_{x_i}^R$ ، اعداد تصادفی تولید شده از توزیع یکنواخت روی بازه $(0, 0.001)$ می‌باشند. اکنون، اعداد فازی مثلی در جدول ۵ [۶] را با پهنای متفاوت تولید شده در جدول ۲ گردآوری می‌کنیم. در نتیجه، می‌خواهیم شاخص‌های تعمیم یافته بر اساس قضیه‌های ۵.۳ و ۶.۳ را با فرض اعداد فازی مثلی نامتقارن شبیه سازی شده در جدول ۲ و کران‌های فازی خطی و نمایی محاسبه کنیم. همچنین، $m = 1$ برای تابع وزنی g در تعریف ۳.۲ و ۲.۳ در نظر می‌گیریم. به عنوان نمونه، روند محاسبات برای برآورد شاخص کارایی تعمیم یافته $C_{\tilde{p}}$ به ازای داده مثلی و حدود مشخصه خطی، با توجه به فرمول (۷.۳) به شرح ذیل است

$$\begin{aligned}\hat{C}_{\tilde{p}} &= \frac{D(\widetilde{USL}_L, \widetilde{LSL}_L)}{6\sqrt{v_T(\tilde{x})}} \\ &= \frac{1}{m+2}[(m+1)(u_1 - l_1) + (u_0 - l_0)] \\ &\quad \div 6\left[v(x) + \frac{1}{(m+2)(m+3)}[v(s_x^R) + v(s_x^L)]\right. \\ &\quad \left.+ \frac{1}{m+2}[\text{cov}(x, s_x^R) - \text{cov}(x, s_x^L)]\right]^{\frac{1}{2}} \\ &= 0.482481624\end{aligned}$$

همان‌طور که در جدول ۱ مشاهده می‌کنید، فرایند قطر پیستون کارا نیست، چون که مقادیر برآورد شده از $C_{\tilde{p}}$ و $C_{\tilde{pm}}$ خیلی کوچکتر از مقدار بحرانی ۱.۳۳ می‌باشند. مهندسین کیفیت باید این نتیجه را در روند پیشرفت کارایی فرایند در نظر بگیرند.

جدول ۱: برآوردی از شاخص‌های کارایی تعمیم یافته با در نظر گرفتن اندازه‌های فازی گزارش شده با پهنای متفاوت

خطی	نمایی
$\widetilde{USL}_L = (74.0146, 74.0168)_{UL}$	$\widetilde{USL}_E = (74.0146, 0.0022)_{UE}$
$\widetilde{LSL}_L = (73.9856, 73.9878)_{LL}$	$\widetilde{LSL}_E = (73.9878, 0.0022)_{LE}$
$T = 74.0012$	$T = 74.0012$
$\hat{C}_{\tilde{p}} = 0.482481624$	$\hat{C}_{\tilde{p}} = 0.504328523$
$\hat{C}_{\tilde{pm}} = 0.482083247$	$\hat{C}_{\tilde{pm}} = 0.504295098$

جدول ۲: اعداد فازی مثلثی نامتقارن تولید شده با پهناهای متفاوت

ردیف	داده‌های مثلثی قطر رینگ پیستون
۱	$T(74.0025, 4.39 \times 10^{-4}, 9.24 \times 10^{-4})$
۲	$T(74.0065, 0.417 \times 10^{-4}, 1.2 \times 10^{-4})$
۳	$T(74.0085, 6.37 \times 10^{-4}, 1.98 \times 10^{-4})$
۴	$T(73.9905, 5.8 \times 10^{-4}, 5.34 \times 10^{-4})$
۵	$T(74.0145, 9.75 \times 10^{-4}, 4.47 \times 10^{-4})$
۶	$T(73.9865, 0.347 \times 10^{-4}, 5.85 \times 10^{-4})$
۷	$T(73.9985, 5.27 \times 10^{-4}, 8.4 \times 10^{-4})$
⋮	⋮
۱۲۵	$T(74.0105, 4.54 \times 10^{-4}, 1.66 \times 10^{-4})$

۴. دست‌آوردهای پژوهش

از شاخص‌های کارایی فرایند می‌توان به عنوان ابزاری مفید و تأثیرگذار برای اندازه‌گیری و میزان سنجش کیفیت محصول و عملکرد فرایند بهره‌برد. این شاخص‌ها ابزارهای آماری مفیدی برای خلاصه‌سازی پراکندگی و وضعیت فرایند با استفاده از تجزیه و تحلیل کارایی فرایند هستند. با این حال، برخی از محدودیت‌ها همانند تعریف غیرفازی از کیفیت و تعریف دقیق از داده‌ها وجود دارند که مانع تجزیه و تحلیل عمیق و انعطاف‌پذیر می‌شوند. در این مقاله، یک نسخه تعمیم یافته از شاخص‌های کارایی ژوران و تاگوچی برای برآورد کارایی فرایند بر اساس کیفیت فازی و داده‌های فازی-مقدار نامتقارن پیشنهاد شده است. این شاخص‌های تعمیم یافته قادر به ارزیابی فرایندهای فازی توسط ترکیب برداری از اعداد فازی، یک مقدار هدف دقیق و دو تابع عضویت از حدود مشخصه فنی فازی هستند. برآوردهای فاصله‌ای از شاخص‌های کارایی تعمیم یافته و آزمون کارایی فرایندهای فازی دو مسئله جالب هستند که می‌توان در کارهای آینده در نظر گرفت.

مراجع

۱. ع. پرچمی، م. ماشینچی، کیفیت فازی و نسل جدیدی از شاخص‌های کارایی، اندیشه‌آماري ۱۳ (۱۳۸۶)، شماره ۱ ۷۶-۶۸.
۲. ع. پرچمی، آزمون‌های فرض و شاخص‌های کارایی فرایند در محیط فازی، رساله دکتری آمار، دانشگاه فردوسی مشهد، ۱۳۹۴.
3. Chen, T.W., Lin, J.Y. and Chen, K.S. (2003), *Selecting a supplier by fuzzy evaluation of capability indices cpm*, The International Journal of Advanced Manufacturing Technology 22, 534-540
4. Hsiang, T.C. (1985), *A tutorial on quality control and assurance-the taguchi methods*, ASA Annual Meeting, Las Vegas Nevada, USA.
5. Juran, J.M. and Gryna F.M. (1988), *Juran,s quality control handbook*, mc-grawhill, New York, 24-26.
6. Kaya, I. and Kahraman, C. (2011), *Process capability analyses based on fuzzy measurements and fuzzy control charts*, Expert Systems with Applications 38, 3172-3184.
7. Lee, H.T. (2001), *Cpk index estimation using fuzzy numbers*, 683-688.

8. Parchami, A. and Mashinchi, M. (2010), *A new generation of process capability indices*, Journal of Applied Statistics 37, 77–89.
9. Parchami, A. and Mashinchi, M. (2007), *Fuzzy estimation for process capability indices*, Information Sciences 177, 1452–1462.
10. Parchami, A., Sadeghpour-Gildeh, B., Nourbakhsh, M., and Mashinchi, M. (2014) *A new generation of process capability indices based on fuzzy measurements*, Journal of Applied Statistics 41, 1122–1136.



نمودار کنترل کیفیت فازی- آماری فرایند برای جمعیت چوله

مهسا صعصعی^{1*}، ماشالله ماشین چی²، رضا پور موسی³، و فرزانه هاشمی⁴

^{1,2,3,4} گروه آمار، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان

Mahsasasaei@gmail.com

mashinchi@yahoo.com

pourmusa@yahoo.com

farzane.hashemi1367@yahoo.com

چکیده. نظریه مجموعه های فازی و نظریه احتمال همواره برای مدل سازی موارد مربوط به عدم قطعیت مورد استفاده قرار می گیرند. عدم قطعیت به دلیل عدم وجود اطلاعات کامل و دقیق و عدم قطعیت به دلیل ماهیت تصادفی ذاتی فرایندها، اما اکثر مواقع در سیستم های واقعی با مواردی مواجه هستیم که فرایندها تحت تأثیر هر دو نوع عدم قطعیت قرار دارند. این مقاله به مدل سازی یک نمودار کنترل فازی- آماری فرایند در حضور این دو نوع عدم قطعیت برای جمعیت های چوله می پردازد و با استفاده از یک مثال صنعتی، این موضوع را به تصویر می کشد.

۱. پیش گفتار

پربکارترین ابزارهای کنترل آماری فرایند، استفاده از نمودار کنترل کیفی می باشد. در حالت استاندارد فرض می شود که مشخصه کیفیت بصورت نرمال توزیع شده باشد، اما در سیستم های واقعی همواره فرضیه نرمال بودن برقرار نیست و با مواردی مواجه هستیم که فرایندها چوله باشند مانند فرایندهای شیمیایی، فرایندهای نیمه هادی و فرایندهای سایش ابزار برش. در نمودارهای کنترل کیفی استاندارد برای جمعیت های چوله، نرخ هشدار خطا با افزایش چولگی زیاد می شود. بنابراین نیازمند شیوه ای هستیم که بتواند نرخ هشدار خطا را در سطح مطلوبی کنترل کند. در مواجهه با این نوع فرایندها سه شیوه رایج است: [۲]

۱- یافتن تبدیل مناسب نزدیک به مدل نرمال

۲- افزایش حجم نمونه و استفاده از قضیه حد مرکزی

۳- نمودارهای کنترل کیفی تجربی

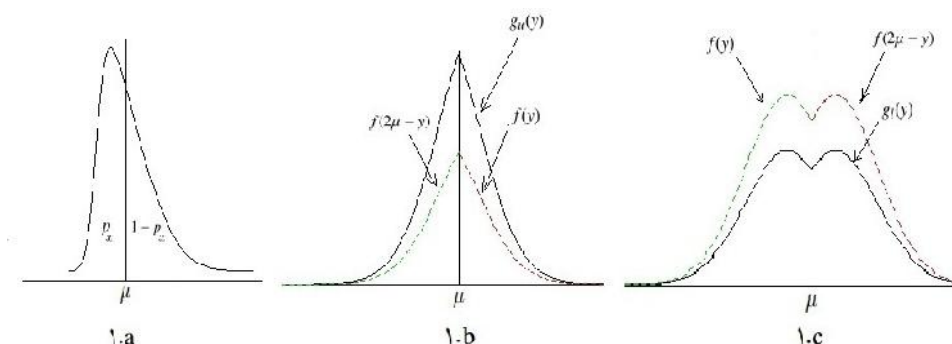
در بخش بعدی این مقاله، ابتدا نمودارهای کنترل کیفی تجربی را بر پایه انحراف معیار موزون شرح داده ایم. در بخش سوم، فرض بر آن است که داده ها و اطلاعات دقیق و قطعی است، لذا به توصیف نمودارهای کنترل آماری به شیوه مشروح در بخش دوم برای فرایندهای چوله می پردازیم.

واژگان کلیدی. حدود کنترل کیفی، جمعیت چوله، انحراف استاندارد وزنی، فرایند فازی- آماری.
* سخنران

سپس در بخش چهارم با فرض اینکه فرایندهای چوله تحت تأثیر هر دو نوع عدم قطعیت ناشی از عدم وجود اطلاعات کامل و دقیق (فازی) و ماهیت تصادفی ذاتی فرایندها (آماري) قرار دارند، عملکرد نمودارهای کنترل فازی-آماري را مورد بررسی قرار داده و در بخش پایانی آن در قالب یک مثال عددی، کاربرد روش مذکور نشان داده شده است.

۲. روش تجربی بر پایه انحراف معیار موزون

در این بخش یک روش تجربی بر پایه انحراف معیار موزون (WSD)، برای تقریب توزیع چوله بوسیله توزیع نرمال شرح می دهیم [۲]. این روش بر این اساس بنا شده که هر توزیع چوله را می توان به دو بخش روی نقطه میانگین تقسیم کرد به گونه ای که هر کدام جهت ساختن یک توزیع متقارن جدید مورد استفاده قرار گیرد. شکل ۱ ایده اصلی این روش را نشان میدهد. فرض کنید f تابع چگالی احتمال مشخصه کیفیت X با میانگین μ و انحراف معیار σ باشد و $P_x = Pr(X \leq \mu)$ (شکل ۱.a). قرینه سمت راست f را نسبت به خط عمودی در نقطه μ رسم کرده سپس تابع چگالی احتمال متقارن g_U را طوری که در خواص تابع چگالی صدق کند اصلاح می کنیم (شکل ۱.b). بطور مشابه، تابع چگالی متقارن g_L نیز از قرینه سمت چپ نسبت به خط عمودی در نقطه μ حاصل می شود. (شکل ۱.c) تابع چگالی های جدید g_U و g_L دارای میانگین های یکسان μ و به ترتیب انحراف معیارهای



شکل ۱: (۱.a) تابع چگالی چوله، (۱.b) و (۱.c) تابع چگالی های متقارن متناظر با سمت راست و چپ.

متفاوت σ_U و σ_L می باشند. بنابراین

$$g_U(y) = \begin{cases} \frac{1}{2(1-P_x)} f(2\mu - y) & \text{if } y \leq \mu, \\ \frac{1}{2(1-P_x)} f(y) & \text{if } y > \mu. \end{cases} \quad (1.2)$$

$$g_L(y) = \begin{cases} \frac{1}{2P_x} f(y) & \text{if } y \leq \mu, \\ \frac{1}{2P_x} f(2\mu - y) & \text{if } y > \mu. \end{cases} \quad (2.2)$$

که در آن $\frac{1}{2(1-P_x)}$ و $\frac{1}{2P_x}$ ثابت های وزنی هستند به طوری که g_U و g_L در خواص تابع چگالی صدق کنند. از رابطه (۱,۲) و (۲,۲) داریم:

$$\sigma_U = \sqrt{\frac{1}{1 - P_x} \int_{\mu}^{\infty} (y - \mu)^2 f(y) dy}, \quad (3.2)$$

$$\sigma_L = \sqrt{\frac{1}{P_x} \int_{-\infty}^{\mu} (y - \mu)^2 f(y) dy}. \quad (4.2)$$

از طرفی با استفاده از تقریب نیم واریانس چوبینه و برنتینگ [۳] داریم:

$$\int_{\mu}^{\infty} (y - \mu)^2 f(y) dy \cong P_x \sigma^2, \quad (5.2)$$

$$\int_{-\infty}^{\mu} (y - \mu)^2 f(y) dy \cong (1 - P_x) \sigma^2. \quad (6.2)$$

بنابراین تقریب های σ_U و σ_L بصورت زیر بدست می آیند:

$$\sigma_U \cong \sqrt{\frac{P_x}{1 - P_x}} \sigma, \quad \sigma_L \cong \sqrt{\frac{1 - P_x}{P_x}} \sigma. \quad (7.2)$$

با توجه به روابط زیر

$$\frac{\sigma_U^W}{\sigma} = \frac{2\sigma_U^W}{2\sigma_U^W + 2\sigma_L^W} = \frac{\sigma_U}{\sigma_U + \sigma_L} = P_x \quad (8.2)$$

$$\frac{\sigma_L^W}{\sigma} = \frac{2\sigma_L^W}{2\sigma_U^W + 2\sigma_L^W} = \frac{\sigma_L}{\sigma_U + \sigma_L} = (1 - P_x) \quad (9.2)$$

و با استفاده از رابطه (۷،۲)، انحراف معیارهای وزنی بصورت تقریبی برابر است با

$$\sigma_U^W \cong P_x \sigma, \quad \sigma_L^W \cong (1 - P_x) \sigma \quad (10.2)$$

و در نهایت برای هر یک از توابع چگالی احتمال متقارن g_L و g_U ، تقریب های تابع چگالی احتمال نرمال f_L و f_U به ترتیب به فرم زیر ساخته می شوند:

$$f_L(x) = \frac{1}{2\sigma_L^W} \phi\left(\frac{x - \mu}{2\sigma_L^W}\right), \quad (11.2)$$

$$f_U(x) = \frac{1}{2\sigma_U^W} \phi\left(\frac{x - \mu}{2\sigma_U^W}\right) \quad (12.2)$$

که در آن $\phi(\cdot)$ تابع چگالی احتمال نرمال استاندارد می باشد.

توجه ۱۰.۲. اگر توزیع تحت بررسی متقارن باشد، آنگاه $P_x = \frac{1}{2}$ و $\sigma_U^W = \sigma_L^W = \frac{1}{2}\sigma$. در صورتی که چوله به راست باشد، $P_x > \frac{1}{2}$ و $\sigma_U^W > \sigma_L^W$ و در صورتی که چوله به چپ باشد، $P_x < \frac{1}{2}$ و $\sigma_U^W < \sigma_L^W$.

۳. نمودار کنترلی میانگین بر اساس روش انحراف معیار موزون در حالت استاندارد فرض کنید مشخصه کیفیت دارای توزیع $X \sim N(\mu, \sigma^2)$ باشد. حدود کنترل نمودارهای کنترلی $\bar{X}$ بصورت زیر می باشند:

$$\begin{aligned} UCL_{STD} &= \mu + Z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \\ LCL_{STD} &= \mu - Z_{\alpha/2} \frac{\sigma}{\sqrt{n}} \end{aligned} \quad (۱.۳)$$

که در آن $Z_{\alpha/2}$ چندک $100(1 - \alpha/2)\%$ ام توزیع نرمال استاندارد است. در صورتی که پارامترها نامعلوم باشند، آنگاه میانگین و انحراف معیار بوسیله $\bar{X}$ و $\frac{\bar{R}}{d_2}$ برآورد می شوند که در آن $\bar{R}$ میانگین دامنه ای و d_2 ثابت نمودار کنترلی توزیع نرمال می باشد.

توجه ۱.۳. معمولاً در نمودارهای شوهارت 3σ عدد ۳ جایگزین $Z_{\alpha/2}$ می شود.

توجه ۲.۳. به خاطر عملکرد ضعیف حدود کنترل استاندارد در نمودارهای کنترلی فرایندهای چوله از روش WSD معرفی شده در بخش ۲ استفاده می کنیم.

در روش WSD با استفاده از تقریب های نرمال در حدود بالایی و پایینی حدود کنترلی استاندارد به جای σ از $2\sigma_U^W$ و $2\sigma_L^W$ استفاده کرده و با جایگذاری رابطه $(۱, ۲)$ حدود کنترلی بصورت زیر بدست می آیند.

قضیه ۳.۳. فرض کنید مشخصه کیفیت X متغیر تصادفی چوله از توزیعی با پارامترهای معلوم باشد، آنگاه حدود کنترل بصورت زیر می باشند: [۲]

$$\begin{aligned} UCL_{WSD} &= \mu + Z_{\alpha/2} \frac{\sigma}{\sqrt{n}} 2P_x \\ LCL_{WSD} &= \mu - Z_{\alpha/2} \frac{\sigma}{\sqrt{n}} 2(1 - P_x) \end{aligned} \quad (۲.۳)$$

قضیه ۴.۳. فرض کنید مشخصه کیفیت X متغیر تصادفی چوله از توزیعی با پارامترهای نامعلوم باشد، آنگاه حدود کنترل بصورت زیر می باشند: [۲]

$$\begin{aligned} UCL_{STD} &= \hat{\mu} + Z_{\alpha/2} \frac{\hat{\sigma}}{\sqrt{n}} 2\hat{P}_x \\ LCL_{STD} &= \hat{\mu} - Z_{\alpha/2} \frac{\hat{\sigma}}{\sqrt{n}} 2(1 - \hat{P}_x) \end{aligned} \quad (۳.۳)$$

که در آن برآوردهای نااریب پارامترهای μ ، σ و P_x به قرار زیر است:

$$\hat{\mu} = \bar{\bar{X}}, \quad \bar{\bar{X}} = \frac{\sum_{i=1}^k \sum_{j=1}^n X_{ij}}{k \times n}$$

$$\hat{\sigma} = \frac{\bar{R}}{d_2^{WSD}}, \quad d_2^{WSD} \equiv \hat{P}_x d_2(2n(1 - \hat{P}_x)) + (1 - \hat{P}_x) d_2(2n\hat{P}_x)$$

$$\hat{P}_x = \frac{\sum_{i=1}^k \sum_{j=1}^n I(\bar{\bar{X}} - X_{ij})}{k \times n}, \quad I(x) = \begin{cases} 1 & \text{if } x \geq 0, \\ 0 & \text{if o.w.} \end{cases}$$

بطوری که $\bar{\bar{X}}$ میانگین k نمونه n تایی از مشاهدات و d_2^{WSD} ثابت نمودار کنترلی توزیع چوله به روش انحراف معیار موزون است.

توجه ۵.۳. اگر توزیع تحت بررسی متقارن باشد روابط $(2, 3)$ و $(3, 3)$ به حالت استاندارد کاهش می یابد. اگر توزیع چوله به راست باشد فاصله UCL از خط مرکزی CL بزرگتر از LCL است و اگر توزیع چوله به چپ باشد فاصله LCL از CL بزرگتر از UCL است.

۴. نمودار کنترل فازی-آماري

در عمل موقعیت های زیادی وجود دارند که فرایندها دستخوش دو نوع عدم قطعیت فازی و آماری هستند. به عنوان مثال فرایندی را با یک مشخصه کیفی در نظر بگیرید که هدف آن کنترل مشخصه کیفی تصادفی در حضور عدم قطعیت حاکی از دستگاههای اندازه گیری و یا قضاوت افراد است. نمودارهای کنترل آماری از آنجا که عدم قطعیت فازی را در نظر نمی گیرند مناسب نمی باشند. از طرفی، نمودارهای کنترل فازی نیز ماهیت تصادفی بودن مشخصه کیفی را نادیده می گیرند. لذا بناسازی نمودارهایی که قادر به کنترل عدم قطعیت های فازی و آماری بطور همزمان باشند، از اهمیت خاصی برخوردار است که در این بخش به آن می پردازیم.

۴-۱. ناحیه کنترل فازی و متغیرهای تصادفی فازی

در این مقاله از مفهوم متغیرهای تصادفی فازی در مدل کردن فرایندها در حضور دو عدم قطعیت مذکور استفاده می کنیم. در این راستا، مشخصه کیفی فرایند تحت تأثیر دو تابع قرار دارد. تابع توزیع احتمال مشخصه کیفی، تغییرپذیری ذاتی فرایند را مدل می کند و ابهام موجود در فرایند به وسیله یک تابع عضویت فازی مدل سازی می شود. شاپیرو [۶]، بطور جامع مفاهیم و تعاریف مختلف متغیرهای تصادفی فازی را مورد بررسی قرار داده است. فرض کنید که مشخصه کیفی فرایند X دارای توزیع چوله f_θ است. از طرفی به دلیل حضور عدم قطعیت فازی، مشاهدات در قالب یک عدد فازی LR به فرم $\tilde{X}_{LR} = \langle m, s, l, r \rangle_{LR}$ گزارش داده می شوند که در آن m ، مرکز عدد فازی را نشان می دهد. s بیانگر نصف طول هسته و r و l به ترتیب بیانگر پهنای

راست و چپ عدد فازی می باشند. تابع عضویت این عدد فازی به فرم زیر است:

$$M_{\tilde{X}}(x) = \begin{cases} 0 & \text{if } x < m - s - l, \\ L\left(\frac{m-s-x}{l}\right) & \text{if } m - s - l \leq x < m - s \\ 1 & \text{if } m - s \leq x < m + s \\ R\left(\frac{x-s-m}{r}\right) & \text{if } m + s \leq x < m + s + r \\ 0 & \text{if } x > m + s + r \end{cases} \quad (۱.۴)$$

که در آن L و R توابعی از چپ پیوسته و غیر صعودی از $[0, 1] \rightarrow [0, 1]$ هستند بطوریکه $L(1) = R(1) = 0$ و $L(0) = R(0) = 1$.

با فرض اینکه مشخصه کیفی فرایند، متغیر تصادفی فازی دارای توزیع چوله با میانگین فازی $\tilde{\mu}_{LR} = \langle \mu_m, \mu_s, \mu_l, \mu_r \rangle_{LR}$ و واریانس دقیق σ^2 و همچنین فرض مجهول بودن پارامترها ابتدا لازم است که این دو پارامتر بصورت زیر تخمین زده شوند.

نکته ۱.۴. زمانی که فرض می شود فرایند تحت کنترل است، نمونه هایی به حجم n و به تعداد k از فرایند جمع آوری می شوند و بوسیله آنالیز این نمونه های اولیه، ابتدا میانگین فرایند به فرم زیر تخمین زده می شود:

$$\begin{aligned} \hat{\mu}_{LR} &= \langle \hat{\mu}_m, \hat{\mu}_s, \hat{\mu}_l, \hat{\mu}_r \rangle_{LR} = \frac{\sum_{i=1}^k \sum_{j=1}^n \langle m_{ij}, s_{ij}, l_{ij}, r_{ij} \rangle_{LR}}{kn} \\ &= \frac{\sum_{i=1}^k \langle m_i, s_i, l_i, r_i \rangle_{LR}}{k} = \left\langle \frac{\sum_{i=1}^k \bar{m}_i}{k}, \frac{\sum_{i=1}^k \bar{s}_i}{k}, \frac{\sum_{i=1}^k \bar{l}_i}{k}, \frac{\sum_{i=1}^k \bar{r}_i}{k} \right\rangle_{LR} \\ &= \langle \bar{\bar{m}}, \bar{\bar{s}}, \bar{\bar{l}}, \bar{\bar{r}} \rangle_{LR} \end{aligned}$$

نکته ۲.۴. کورنر [۵]، نشان داد که واریانس متغیر تصادفی فازی به فرم کلی زیر قابل محاسبه است:

$$Var(\tilde{X}_{LR}) = \sigma_m^2 + \sigma_s^2 + \frac{\sigma_l^2}{2} \left(\int_0^1 (L^{(-1)(\alpha)})^2 d\alpha \right) + \frac{\sigma_r^2}{2} \left(\int_0^1 (R^{(-1)(\alpha)})^2 d\alpha \right)$$

زمانی که مشاهدات فازی به فرم اعداد دوزنقه ای باشند، آن گاه می توان نوشت:

$$\sigma^2 = Var(\tilde{X}_{LR}) = \sigma_m^2 + \sigma_s^2 + \frac{\sigma_l^2}{6} + \frac{\sigma_r^2}{6}$$

که برآورد نااریب آن به قرار زیر است:

$$\bar{s}^2 = Var(\tilde{X}_{LR}) = \bar{s}_m^2 + \bar{s}_s^2 + \frac{\bar{s}_l^2}{6} + \frac{\bar{s}_r^2}{6}.$$

شایان ذکر است، برآورد گر نااریب انحراف استاندارد به فرم $\hat{\sigma} = \frac{\sqrt{\bar{s}^2}}{c_4}$ است که در آن

$$c_4 = \sqrt{\frac{2}{n-1}} \times \frac{\Gamma(\frac{n}{2})}{\Gamma(\frac{n-1}{2})}$$

قضیه ۳.۴. فرض کنید مشخصه کیفیت فازی $\tilde{X}_{LR}$ متغیر تصادفی چوله از توزیعی با پارامترهای نامعلوم باشد، حدود کنترل فازی در سطح خطای نوع اول α به صورت زیر می باشند:

$$\begin{aligned} U\tilde{C}L_{LR} &= \hat{\mu}_{LR} + Z_{\alpha/2} \frac{\hat{\sigma}}{\sqrt{n}} 2\hat{P}_x = \langle \bar{m}, \bar{s}, \bar{l}, \bar{r} \rangle_{LR} + Z_{\alpha/2} \frac{\hat{\sigma}}{\sqrt{n}} 2\hat{P}_x \\ &= \langle \bar{m} + Z_{\alpha/2} \frac{2\hat{P}_x \hat{\sigma}}{\sqrt{n}}, \bar{s} + Z_{\alpha/2} \frac{2\hat{P}_x \hat{\sigma}}{\sqrt{n}}, \bar{l} + Z_{\alpha/2} \frac{2\hat{P}_x \hat{\sigma}}{\sqrt{n}}, \bar{r} + Z_{\alpha/2} \frac{2\hat{P}_x \hat{\sigma}}{\sqrt{n}} \rangle_{LR} \\ \tilde{C}L_{LR} &= \hat{\mu}_{LR} = \langle \bar{m}, \bar{s}, \bar{l}, \bar{r} \rangle_{LR} \\ L\tilde{C}L_{LR} &= \hat{\mu}_{LR} - Z_{\alpha/2} \frac{\hat{\sigma}}{\sqrt{n}} 2(1 - \hat{P}_x) = \langle \bar{m}, \bar{s}, \bar{l}, \bar{r} \rangle_{LR} - Z_{\alpha/2} \frac{\hat{\sigma}}{\sqrt{n}} 2(1 - \hat{P}_x) \\ &= \langle \bar{m} - Z_{\alpha/2} \frac{2(1 - \hat{P}_x) \hat{\sigma}}{\sqrt{n}}, \bar{s} - Z_{\alpha/2} \frac{2(1 - \hat{P}_x) \hat{\sigma}}{\sqrt{n}}, \bar{l} - Z_{\alpha/2} \frac{2(1 - \hat{P}_x) \hat{\sigma}}{\sqrt{n}}, \bar{r} - Z_{\alpha/2} \frac{2(1 - \hat{P}_x) \hat{\sigma}}{\sqrt{n}} \rangle_{LR} \end{aligned} \quad (2.4)$$

که در آن $\hat{\sigma}$ برآورد انحراف معیار متغیر تصادفی فازی است و $\hat{P}_x = \frac{\sum_{i=1}^k \sum_{j=1}^n I(\hat{\mu}_{LR} - \tilde{X}_{ijLR})}{k \times n}$ برهان. با استفاده از قضیه ۴.۳، نکات ۱.۴ و ۲.۴ حدود کنترل فازی قابل محاسبه می باشند. $\square$

۲-۴. آماره نمودار کنترل فازی-آماري

تعریف. دو مجموعه فازی، نرمال، محدب و حقیقی F_1 و F_2 را در مجموعه مرجع X در نظر بگیرید. $\mu_{F_1}(x)$ و $\mu_{F_2}(x)$ به ترتیب بیانگر تابع عضویت دو مجموعه فازی می باشند. درصدی از مجموعه فازی F_1 که خارج مجموعه فازی F_2 قرار می گیرد به فرم زیر قابل محاسبه است: [۴]، [۱]

$$A_{F_1}^{out} = \frac{\int \min(\mu_{F_1}(x), \mu_{F_2}(x)) dx}{\int \mu_{F_1}(x) dx} \quad (3.4)$$

عدد فازی دوزنقه ای F_1 و ناحیه کنترل فازی دوزنقه ای F_2 را در نظر بگیرید، آنگاه تصمیم گیری در مورد فرایند در قالب تابع زیر خلاصه شده است. [۱]

$$\text{process state} = \begin{cases} \text{in control} & A_{sample}^{out} = 0, \\ \text{rather in control} & 0 < A_{sample}^{out} \leq 0.2, \\ \text{rather out of control} & 0.2 < A_{sample}^{out} \leq 0.5, \\ \text{out of control} & 0.5 < A_{sample}^{out} \leq 1. \end{cases}$$

۳-۴. مثال کاربردی

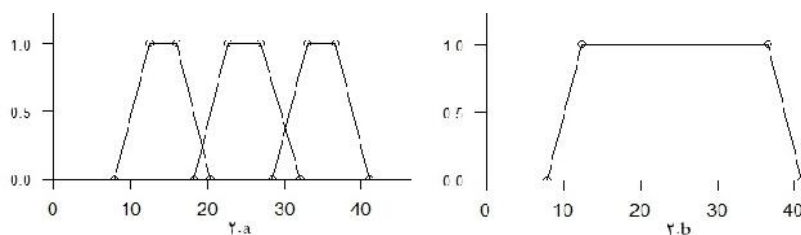
مثال. کارخانه سازنده وسایل الکترونیکی خواهان کنترل کیفیت خازن های تولیدی است. مقاومت الکتریکی خازن یکی از مشخصات مهم کیفی است که با استفاده از دستگاه های حساس الکتریکی سنجیده می شوند. از طرفی با در نظر گرفتن خطای اندازه گیری و همچنین شرایط محیطی آزمایشگاه، داده ها در قالب اعداد فازی دوزنقه ای گزارش می شوند [۱]. با

آنالیز اولیه داده ها و $\alpha = 0.0027$ ، ناحیه کنترل فازی در قالب عدد فازی دوزنقه ای (۷/۸۴، ۱۲/۳۷، ۳۶/۵۴، ۴۰/۹۹) محاسبه می گردد. جدول ۱، بیانگر میانگین ۳۰ نمونه متوالی

جدول ۱: داده های فازی مقاومت الکتریکی خازن

i	a	c	b	d	A_{STD}^{out}	A_{WSD}^{out}
۱	۲۵	۳۰	۳۰	۳۵	۰	۰
۲	۱۵	۲۰	۳۰	۳۵	۰	۰
۳	۴	۵	۱۲	۱۵	۰/۰۴۹۵	۰
۴	۳	۶	۶	۸	۰/۰۵۱۶	۰
۵	۳۲	۳۸	۳۸	۴۵	۰	۰
۶	۱۶	۲۰	۲۴	۲۸	۰	۰
۷	۳	۴	۸	۱۲	۰/۱۴۹۳	۰
۸	۲۷	۳۶	۴۴	۵۰	۰/۰۷۹	۰
۹	۹	۱۱	۱۵	۲۰	۰	۰
۱۰	۷	۱۰	۱۳	۱۵	۰	۰
۱۱	۳	۶	۶	۱۰	۰/۰۴۱۲	۰
۱۲	۲۷	۳۲	۳۲	۳۷	۰	۰
۱۳	۱۱	۱۳	۱۳	۱۵	۰	۰
۱۴	۳۹	۵۰	۵۲	۵۵	۰/۶۸۰۳	۰/۳۰۲۵
۱۵	۲۸	۳۸	۴۱	۴۵	۰	۰
۱۶	۳۳	۴۰	۴۰	۴۴	۰	۰
۱۷	۲۸	۳۲	۵۰	۶۰	۰/۳۶۹	۰/۱۵۴۱
۱۸	۳۳	۳۹	۳۹	۴۳	۰	۰
۱۹	۱۲	۱۵	۲۱	۲۸	۰	۰
۲۰	۲۳	۲۸	۲۸	۳۶	۰	۰
۲۱	۲۸	۳۲	۳۵	۴۲	۰	۰
۲۲	۱۴	۱۸	۲۸	۳۳	۰	۰
۲۳	۲۴	۳۰	۳۰	۳۴	۰	۰
۲۴	۲۰	۲۵	۲۵	۳۱	۰	۰
۲۵	۲۵	۳۱	۴۱	۴۶	۰	۰
۲۶	۷	۱۰	۲۵	۲۸	۰	۰
۲۷	۳	۵	۱۴	۲۰	۰/۰۴۲۹	۰
۲۸	۲۳	۲۸	۳۵	۳۸	۰	۰
۲۹	۱۷	۲۰	۲۵	۲۹	۰	۰
۳۰	۵	۸	۸	۱۵	۰	۰

هریک به حجم ۴ خازن است. میزان خروج هر یک از نمونه های فازی از ناحیه تحت کنترل فازی توسط رابطه (۳،۴) محاسبه، در ستون آخر جدول ۱ آورده شده است و با میزان خروج از ناحیه تحت کنترل فازی استاندارد متناظر در ستون قبلی قابل مقایسه می باشد. شکل (۲.ا)،



شکل ۲: (۲.ا) حدود کنترل فازی بالا و پایین به همراه عدد فازی مرکزی و (۲.ب) ناحیه کنترل فازی متناظر.

حدود کنترل فازی بالایی، مرکزی و پایینی به دست آمده از معادله (۲،۴) و شکل (۲.ب) ناحیه

کنترل فازی را به نمایش می گذارد. تمامی محاسبات و نتایج حاصل از طریق برنامه نویسی در نرم افزار R انجام شده است.

مراجع

۱. ع. فراز، بناسازی نمودار کنترل آماری-فازی فرآیند. مجله ریاضیات کاربردی واحد لاهیجان، سال هشتم، شماره ۲، ۲۹ (۱۳۹۰)، ۴۵-۵۴.
2. Chang, Y. S., Bai, D. S, Control charts for positively-skewed populations with weighted standard deviations. *Quality and Reliability Engineering International*, 17 (2001), 397 –406.
3. Choobineh F., Branting D., A simple approximation for semivariace. *European Journal of Operations Research*, 27 (1986), 364–370.
4. Faraz, A., Shapiro, A. F., An application of fuzzy random variables to control charts. *Fuzzy Sets and Systems*, 161 (2010), 2684 – 2694.
5. Körner, R., On the variance of fuzzy random variables. *Fuzzy Sets and Systems* 92 (1997), 83–93.
6. Shapiro, A. F., Fuzzy random variables. *Insurance: Mathematics and Economics* 44 (2009), 307-314.



نقش خودتنظیمی یادگیری در عملکرد تحصیلی فراگیرندگان: فرا رگرسیون فازی

سید محمود طاهری^۱، اصغر شیرعلی پور^۲، و مسعود اسدی^{۳*}

^۱عضو هیأت علمی دانشکده فنی، دانشگاه تهران، ایران

^۲کارشناس ارشد تحقیقات آموزشی، دانشکده روان‌شناسی و علوم تربیتی، دانشگاه خوارزمی، ایران

asgarshiralipoor@gmail.com

^۳دکتری مشاوره، دانشکده روان‌شناسی و علوم تربیتی، دانشگاه خوارزمی، ایران

چکیده. فرا رگرسیون در مطالعات مرور کمی کاربرد فراوان دارد. از آنجا که مطالعات مرور، جزء مطالعات ثانویه محسوب می‌گردد و در این مطالعات دسترسی به داده‌های پژوهش غالباً ممکن نیست و در نتیجه ابهامات مختلفی، از جمله ابهامات در جمع‌آوری داده‌ها و نیز ابهام در تبدیل آن‌ها به یک مقیاس مشترک (اندازه اثر) وجود دارد، لذا فرا رگرسیون فازی می‌تواند در این زمینه راهگشا و مفید باشد. پژوهش حاضر با معرفی و بهره‌گیری از شیوه فراتحلیل فازی، به تحلیل و ترکیب نتایج پژوهش‌های انجام شده در ایران، در زمینه نقش خودتنظیمی در عملکرد تحصیلی فراگیرندگان پرداخته است. بدین منظور، تعداد ۵۰ پژوهش در زمینه نقش خودتنظیمی در عملکرد تحصیلی، گردآوری و از میان آن‌ها تعداد ۲۷ پژوهش که قابلیت بررسی داشت، برای فراتحلیل انتخاب شد. نتایج پژوهش، با بهره‌گیری از روش ترکیب اندازه اثر با رویکرد فازی (مشاهدات دقیق و فرضیه فازی) نشان داد که متغیر خودتنظیمی به تنهایی حدوداً ۷ درصد از تغییرات عملکرد تحصیلی فراگیرندگان را تبیین می‌کند. نتایج فرا رگرسیون فازی بیانگر نقش متغیرهای مختلف به عنوان متغیرهای تعدیل‌کننده بین متغیر خودتنظیمی و عملکرد تحصیلی است. از یافته‌های این مطالعه می‌توان در طرح‌ریزی برنامه‌های آموزش فراگیرندگان و غنی‌سازی خودتنظیمی و یادگیری مؤثر آن‌ها استفاده کرد.

واژگان کلیدی. فرا رگرسیون فازی، اندازه اثر فازی، خودتنظیمی یادگیری، فراتحلیل.

* سخنران

۱. مقدمه

پژوهش‌ها را می‌توان با توجه به چگونگی جمع‌آوری داده‌ها توسط پژوهش‌گر به دو دسته اولیه و ثانویه دسته‌بندی نمود. مطالعات اولیه به مطالعاتی گفته می‌شود که محقق به صورت مستقیم با نمونه در تماس بوده و نتایج را بر اساس مشاهدات خود به دست می‌آورد. به مطالعاتی که محقق به صورت مستقیم به نمونه‌ها و اطلاعات مربوط به آن‌ها دسترسی نداشته و از اطلاعات جمع‌آوری و یا منتشر شده توسط محققین دیگر استفاده نماید، مطالعات ثانویه اطلاق می‌گردد. یکی از مهم‌ترین مطالعات ثانویه، فراتحلیل می‌باشد. فراتحلیل در دو رویکرد کلی مورد بررسی قرار می‌گیرد: فراتحلیل کمی و کیفی. هدف فراتحلیل کیفی توسعه نظریه‌های متوسط و خرد و یا کلان است که از طریق بررسی مطالعات کیفی به دست می‌آیند و اما فراتحلیل کمی حاصل یکپارچه سازی مطالعات منفرد کمی می‌باشد، هومن (۱۳۸۷) و پیگات (۲۰۱۲). در این مقاله هدف بررسی فراتحلیل کمی با استفاده از رگرسیون فازی می‌باشد.

فرا رگرسیون در مطالعات مرور کمی کاربرد دارد. اگر نتایج آزمون‌های همگنی نشان دهند که اگر اندازه اثر مطالعات همگن نیستند، یا به عبارتی ناهمگون هستند؛ این ناهمگنی ممکن است چندین دلیل عملی داشته باشد. از جمله به کار بردن ابزارهای متفاوت در پژوهش‌های منفرد اولیه، حجم نمونه در مطالعات اولیه، سال انتشار مطالعات، جنسیت، محل جغرافیایی و چندین عامل دیگر. این متغیرها، عامل‌های سوم (متغیرهای تعدیل‌کننده) هستند که ممکن هست رابطه مورد نظر را تحت تاثیر قرار دهند. برای بررسی نقش متغیرهای تعدیل‌کننده در رابطه بین متغیرهای تحت بررسی از فرا رگرسیون استفاده می‌گردد. مبانی نظری فرا رگرسیون، کاملاً شبیه رگرسیون معمولی است؛ به عبارتی باید متغیر وابسته که همان اندازه اثر محاسبه شده در تک تک مطالعات می‌باشد را برگزید و متغیر مستقل هم یک یا چند تا از متغیرهای تعدیل‌کننده می‌باشد، بورن اشتاین و همکاران (۲۰۰۹) و استانلی و جارل (۱۹۸۹).

خودتنظیمی‌یادگیری و فرایندهای مشمول در آن به عنوان شاخص پیشرفت تحصیلی فراگیرندگان، همواره مدنظر محققان تعلیم و تربیت بوده است و پژوهش‌های بسیاری در شناسایی پیشایندها

و پیامدهای آن صورت گرفته است. نتایج تحقیقات حاکی از آن است که دانش آموزان با خود تنظیمی بالا، پیشرفت تحصیلی رضایت بخش و مطلوب و همچنین انگیزه و علاقه بالاتری برای ادامه تحصیل دارند. در دهه اخیر پژوهش‌های متعددی که در رابطه با خودتنظیمی و عملکرد تحصیلی فراگیرندگان در کشور ایران انجام شده است، نتایج متفاوتی را گزارش کرده‌اند. در همه پژوهش‌های انجام گرفته برای تحلیل در این زمینه، داده‌ها، دقیق اندازه‌گیری و در قالب فرمول‌های دقیق، تحلیل شده‌اند. اما به دلیل اینکه مفاهیم خودتنظیمی و عملکرد تحصیلی مفاهیمی منعطف و نادقیق هستند و روابط بین آنها نیز نادقیق است لذا برای تحلیل این مفاهیم و روابط بین آنها بهتر است که از روش‌های نرم و از جمله روش‌های مبتنی بر ریاضیات و منطق فازی استفاده نمود. بر این اساس، هدف پژوهش حاضر بررسی نقش خود تنظیمی در عملکرد تحصیلی فراگیرندگان با رویکرد داده دقیق و فرضیه فازی است.

۲. روش پژوهش

روش مورد استفاده در این پژوهش فراتحلیل با استفاده از شیوه کمی سازی اندازه اثر بوده است. تمامی پایان نامه‌ها و طرح‌های پژوهشی و مقالات منتشر شده در ایران که به رابطه بین خودتنظیمی با عملکرد تحصیلی فراگیرندگان در سال‌های ۱۳۸۰ تا ۱۳۹۴ پرداخته بودند، جامعه آماری پژوهش حاضر را تشکیل دادند. روش نمونه‌گیری، هدفمند و از نوع نمونه‌گیری گلوله برفی بود. از ۵۰ پژوهش یافت شده در قالب پایان‌نامه‌های کارشناسی ارشد و رساله‌های دکتری، طرح‌های پژوهشی و مقالات منتشر شده در پایگاه‌های اطلاعاتی sid.ir، magiran.com، noormages.ir، ensani.ir و Irاندو.آ.آ.آ که به موضوع پژوهش حاضر مرتبط بودند، ۲۷ اثر به دلیل داشتن شرایط حضور در مطالعه انتخاب و بررسی شدند. ملاک انتخاب این اثرها روش پیشنهادی تابانه (Thabane) می‌باشد که عبارت است از:

- ۱- مطالعه با حجم نمونه زیاد، ۲- ابزار اندازه‌گیری مناسب و دارای پایایی و روایی لازم، ۳- پژوهش‌های کارشناسی ارشد و بالاتر، ۴- روش نمونه‌گیری مناسب (نمونه‌گیری تصادفی).

طبق روش فوق، اندازه اثر هر یک از پژوهش‌ها جداگانه محاسبه گردید و با برگرداندن آن‌ها به یک ماتریس مشترک (عمومی) و آنگاه ترکیب آن‌ها، تاثیر متوسط این اندازه اثرها به دست آمد.

۳. روش تحلیل داده‌ها

در فراتحلیل می‌توان اندازه اثر آمارهای t ، F ، Z و X^2 را محاسبه کرد. اندازه اثر به روش‌های مختلفی برآورد می‌شود، بدین شکل که فراتحلیل‌گران با داشتن مقادیر میانگین و انحراف استاندارد گروه‌ها، قادر به محاسبه اندازه اثر هستند. رایج‌ترین شاخص‌ها، r و d هستند که غالباً d برای تفاوت‌های گروهی (محاسبه تفاوت میانگین اندازه اثرها) و r برای پژوهش‌های همبستگی به کار می‌روند، هومن (۱۳۸۷).

جدول ۱: تبدیل آماره‌های آزمون‌های مختلف به شاخص d و r

آماره	فرمول شاخص r	آماره	فرمول شاخص d
t	$\sqrt{\frac{t^2}{t^2+df}}$	t	$\frac{2t}{\sqrt{df}}$
F	$\sqrt{\frac{F}{F+dferror}}$	F	$\sqrt{\frac{F(n_1+n_2)}{n_1n_2}}$
X^2	$\frac{z}{\sqrt{n}} \cdot \sqrt{\frac{x^2}{N}}$	r	$\frac{2r}{\sqrt{1-r^2}}$

در مطالعه حاضر با استفاده از مدل کوهن، اندازه اثر کم ($0 \leq r \leq 0.1$)، اندازه اثر متوسط ($0.1 < r \leq 0.3$) و اندازه اثر زیاد ($0.3 < r \leq 1$) در نظر گرفته شد، کوهن (۱۹۷۷).

۱.۳. برآورد اندازه اثر کل (روش مبتنی بر آزمون فرضیه فازی). در این پژوهش برای محاسبه اندازه اثر کل از رویکرد رزنتال و روبین استفاده شد، فیلد (۲۰۰۱). در این رویکرد ابتدا اندازه اثرهای پژوهش‌های منفرد از طریق فرمول $z_{r_i} = \frac{1}{2} \log_e \left(\frac{1+r_i}{1-r_i} \right)$ به z فیشر تبدیل می‌گردد. سپس مقدار متوسط اندازه اثرها از فرمول‌های زیر قابل محاسبه است.

فرض کنید z_{r_i} اندازه اثر مربوط به مطالعه i ام باشد، داریم:

$$\bar{z}_{r_i} = \frac{\sum_{i=1}^k w_i z_{r_i}}{\sum_{i=1}^k w_i} = \frac{\sum_{i=1}^k (n_i - 3) z_{r_i}}{\sum_{i=1}^k (n_i - 3)} \quad w_i = n_i - 3 \quad w_i = \frac{1}{v_i} \quad v_i = \frac{1}{n_i - 3}$$

و در نهایت برای بررسی معناداری مقدار $\bar{z}_{r_i}$ به دست آمده، در آزمون فرضیه صفر $\theta : \cong \circ$ از آماره آزمون زیر استفاده می‌گردد.

$$\tilde{Z} = \frac{\bar{z}_{r_i} - \tilde{H}}{SE(\bar{z}_{r_i})} = \left(\sqrt{\sum_{i=1}^k (n_i - 3) \times \frac{\sum_{i=1}^k (n_i - 3)(z_{r_i} - a)}{\sum_{i=1}^k (n_i - 3)}}, \sqrt{\sum_{i=1}^k (n_i - 3) \times \frac{\sum_{i=1}^k (n_i - 3)z_{r_i}}{\sum_{i=1}^k (n_i - 3)}}, \sqrt{\sum_{i=1}^k (n_i - 3) \times \frac{\sum_{i=1}^k (n_i - 3)(z_{r_i} + a)}{\sum_{i=1}^k (n_i - 3)}} \right)_T$$

و

$$SE(\bar{z}_{r_i}) = \sqrt{\frac{1}{\sum_{i=1}^k w_i}} = \sqrt{\frac{1}{\sum_{i=1}^k (n_i - 3)}}$$

عدد فازی تقریباً صفر، در فرضیه صفر، به صورت $\tilde{H} = (-a, \circ, a)_T$ با تابع عضویت زیر منظور شده است، که مقدار a به نظر کاربر بستگی دارد. در این مطالعه $a = 0.11$ در نظر گرفته شده است.

$$\tilde{H}(\lambda) = \begin{cases} \frac{\lambda+a}{a} & -a < \lambda \leq \circ \\ \frac{a-\lambda}{a} & \circ < \lambda \leq a \\ \circ & \lambda > a, \lambda < -a \end{cases}$$

برای تصمیم‌گیری در مورد رد، یا تأیید فرضیه‌ی مورد مطالعه (بین خودتنظیمی یا دگیری و عملکرد تحصیلی فراگیرندگان رابطه وجود دارد) از رویکرد آزمون فرضیه فازی استفاده می‌کنیم. در این مطالعه از روشی برای آزمون فرضیه بر اساس درجه تأیید و رد فرضیه صفر استفاده می‌کنیم که بدین ترتیب است، طاهری و عارفی (۲۰۰۹): وقتی $Z_1 < Z_2$ درجه رد فرضیه صفر برابر است با $DR = \frac{Z_2}{Z_1 + Z_2}$: درجه رد فرضیه H_0 . به عبارتی فرضیه مخالف (H_1) را با همان درجه یعنی $DR = \frac{Z_2}{Z_1 + Z_2}$ می‌پذیریم، یعنی $DR = \frac{Z_1}{Z_1 + Z_2} = 1 - DA$: درجه پذیرش فرضیه H_0 .

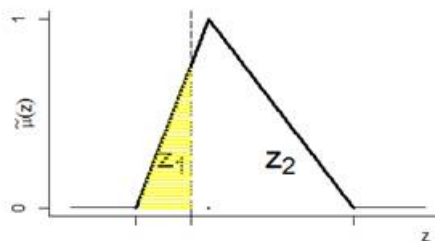
توجه ۱.۳. اگر $DR = ۱$ فرضیه H_0 کاملاً رد می‌شود و فرضیه مخالف کاملاً تأیید می‌گردد. و اگر $DR = ۰$ یعنی فرضیه H_0 کاملاً تأیید می‌شود و فرضیه مخالف کاملاً رد می‌گردد.

وقتی $Z_1 \geq Z_2$ ، درجه پذیرش فرضیه صفر برابر است با

$$DA = \frac{Z_1}{Z_1 + Z_2} \quad \text{درجه پذیرش فرضیه } H_0 \quad DR = \frac{Z_2}{Z_1 + Z_2} = 1 - DA \quad \text{درجه رد فرضیه } H_0$$

به عبارتی فرضیه مخالف (H_1) را با همان درجه یعنی $DR = \frac{Z_2}{Z_1 + Z_2}$ می‌پذیریم، که

$$Z_1 = \int_0^{Z_{1-\alpha}} \tilde{Z}(Z) df, \quad Z_2 = \int_{Z_{1-\alpha}}^0 \tilde{Z}(Z) df$$



شکل ۱: تابع عضویت آماره $\tilde{Z}$ و مقدار بحرانی $Z_{1-\alpha}$

همچنین برای آزمون همگنی بین مطالعات از فرمول زیر بهره گرفته می‌شود، بورن اشتاین و همکاران (۲۰۰۹)

$$\chi^2 = \sum_{i=1}^{i=k} w_i (z_{ri} - \bar{z}_{ri})^2, \quad df = k - 1$$

چنانچه χ^2 محاسبه شده، کوچکتر از χ^2 بحرانی باشد یا به عبارتی χ^2 معنادار باشد ($P > \alpha$) نتیجه می‌گیریم که فرضیه صفر مورد پذیرش قرار می‌گیرد، به عبارتی مطالعات همگن هستند. توجه دارید که فرضیه‌های آزمون همگنی عبارت اند از:

فرضیه صفر: مطالعات همگن هستند، فرضیه مقابل: مطالعات ناهمگن هستند.

۲.۳. بررسی ناهمگنی مطالعات از طریق متغیرهای تعدیل کننده (روش مبتنی بر رگرسیون فازی). برای بررسی علل ناهمگنی بین مطالعات از رگرسیون فازی به خاطر وجود روابط نادقیق بین متغیرهای مورد بررسی استفاده می‌گردد که در زیر به تشریح آن می‌پردازیم، طاهری و ماشین چی (۱۳۸۷).

رگرسیون فازی: در مدل رگرسیون با ضرایب فازی، فرض بر این است که اختلاف بین مقادیر مشاهده شده مربوط به مقادیر وابسته و مقادیر برآورد شده توسط مدل، ناشی از ابهام در ساختار مدل می‌باشد. این ابهام در ضرایب مدل که اعداد فازی هستند منظور می‌شود. بدین ترتیب در این نوع مدل رگرسیون خطی، هدف آن است که بر پایه مجموعه مشاهدات ضرایب فازی طوری بدست آید که: $\tilde{y} = f(\underline{x}) = \tilde{A}_0 + \tilde{A}_1(x_1) + \dots + \tilde{A}_n(x_n)$ یک مدل بهینه باشد. در رگرسیون با ضرایب فازی و مشاهدات غیر فازی (دقیق)، برای بهینه کردن مدل باید نخست اینکه خروجی فازی $\tilde{y}$ ، برای تمامی مقادیر $j = 1, 2, \dots, m$ حداقل داری درجه عضویتی به بزرگی h باشد مقدار h توسط کاربر تعیین می‌گردد و می‌توان آن را به عنوان سطح اعتبار مدل تعبیر کرد. از این شرط، محدودیت‌های مسئله به صورت زیر به دست می‌آید. باید توجه داشت که در نامعادله اخیر از رابطه $s_i^R = k_\circ s_i^L$ استفاده شده است.

$$\begin{aligned} (1-h)k_\circ s_i^L + (1-h) \sum_{i=1}^{i=n} (k_\circ s_i^L x_{ij}) - a_\circ^c - \sum_{i=1}^{i=n} (a_i^c x_{ij}) &\geq -y_j \\ (1-h)s_\circ^L + (1-h) \sum_{i=1}^{i=n} (s_i^L x_{ij}) + a_\circ^c + \sum_{i=1}^{i=n} (a_i^c x_{ij}) &\geq y_j \end{aligned}$$

دوم این که ابهام و یا فازی بودن خروجی مدل، حداقل ممکن باشد. این شرط بیانگر این نقطه است که ابهام در پیش‌بینی متغیر وابسته حداقل باشد. چون بر پایه هر مشاهده، یک خروجی از مدل خواهیم داشت، پس باید مجموع ابهام‌های خروجی‌ها را حداقل کنیم به بیان دیگر باید مقدار زیر حداقل شود

$$\min : Z = (s_o^L + s_o^R) + \sum_{i=1}^{i=n} \left[(s_o^L + s_o^R) \sum_{j=1}^m x_{ji} \right]$$

که تابع هدف فوق را می‌توان بر حسب ضریب کشیدگی $(s_i^R = k_i s_i^L)$ به صورت زیر نوشت

$$\min : Z = m(1 + k_o) s_o^L + \sum_{i=1}^{i=n} \left[(1 + k_i) (s_o^L) \sum_{j=1}^m x_{ji} \right]$$

خاطر نشان می‌شود که در پژوهش حاضر از اعداد فازی مثلثی متقارن بهره گرفته می‌شود. که تابع هدف و محدودیت‌های آن به صورت زیر در نظر گرفته می‌شود

$$\min : Z = \gamma \left(m s_o + \sum_{i=1}^{i=n} \left[(s_i) \sum_{j=1}^m x_{ji} \right] \right)$$

و

$$\begin{aligned} (1-h)s_o + (1-h) \sum_{i=1}^{i=n} (s_i x_{ij}) + a_o^c - \sum_{i=1}^{i=n} (a_i^c x_{ij}) &\geq y_j \\ (1-h)s_o + (1-h) \sum_{i=1}^{i=n} (s_i x_{ij}) - a_o^c + \sum_{i=1}^{i=n} (a_i^c x_{ij}) &\geq -y_j \end{aligned}$$

در نامعادلات بالا h همان برش α و x_{ij} نشان دهنده مشاهده i ام برای متغیر i ام است. می‌توان

تابع هدف را با توجه به γm محدودیت تولید شده توسط m مشاهده مینیم کرد.

یک عدد فازی مانند $\tilde{A}$ را در صورت متقارن بودن به صورت $\tilde{A} = (a, s)$ نیز نشان می‌دهند.

تابع عضویت عدد فازی $\tilde{A}$ تابع دو پارامتر a و s است، که در آن a میانه و s پهنای عدد فازی

است. در واقع عدد فازی $\tilde{A}$ با دو پارامتر a و s نشان دهنده تقریباً a است. پارامترهای

فازی $\tilde{A} = \{\tilde{A}_o, \tilde{A}_1, \dots, \tilde{A}_n\}$ می‌تواند به صورت بردار $\tilde{A} = \{a, s\}$ نشان داده شود، که

$a = (a_o, a_1, \dots, a_n)$ و $s = (s_o, s_1, \dots, s_n)$. بنابر این خروجی رگرسیون به صورت

$\tilde{y} = (a_o, s_o) + (a_1, s_1)x_1 + \dots + (a_n, s_n)x_n$ خواهد بود و لذا تابع عضویت خروجی

رگرسیون به صورت زیر به دست می آید

$$\mu_{\tilde{y}}(y) \begin{cases} 1 - \frac{\left| y - \sum_{i=1}^n a_i x_i \right|}{\sum_{i=1}^n s_i |x_i|} & x \neq c \\ 1 & x, y = \circ \\ \circ & x = \circ, y \neq \circ \end{cases} \quad (1.3)$$

برای ارزیابی مدل برازش شده از شاخص اطمینان $IC = 1 - \frac{SS_E}{SS_T}$ استفاده می شود، **خدایی** (۱۳۸۸)، که $SS_E = 2 \sum_{i=1}^m (y_j - \tilde{y}_i^c)^2$ و $SS_T = \sum_{i=1}^m (y_j - \tilde{y}_i^L)^2 + \sum_{i=1}^m (y_j^R - y_j)^2$ ، مدلی که دارای IC کوچکتری است به عنوان مدل بهینه انتخاب می گردد. ملاک دیگر برای ارزیابی مدل، شاخص MDM یعنی متوسط درجات عضویت، به صورت زیر است، پوراحمد و همکاران (۲۰۱۱)

$$MDM = \frac{1}{m} \sum_{i=1}^m \tilde{Y}_i(y_i), \quad (2.3)$$

۴. کاربرد در خودتنظیمی یادگیری

تحلیل آماری این مطالعه با استفاده از نرم افزار فراتحلیل جامع و R با رویکرد ترکیب اندازه اثر روزنتال و روبین، **فیلد** (۲۰۰۱)، صورت گرفت. اندازه اثر ۲۷ مطالعه در جدول ۲ ارائه شده است. جدول ۲ نشان می دهد که از ۲۷ مطالعه مورد بررسی، در سطح ۵ درصد، پژوهش های با شماره ۷، ۲۱ و ۲۲ معنادار نیستند و سایر پژوهش ها معنادار هستند. بزرگترین اندازه اثر مربوط به مطالعه شماره ۱۴، کوچکترین آن مربوط به مطالعه شماره ۲۲ است.

۱.۴. سوگیری انتشار. بخشی از هر فراتحلیل، ارزیابی سوگیری انتشار است که ناشی از انتشار پژوهش های چاپ شده (معنادار) و عدم انتشار پژوهش های چاپ نشده (غیر معنادار) و انواع خطاها می باشد. هر فراتحلیل به سبب ملاک های انتخاب و حذف پژوهش ها مقداری

جدول ۲: برآورد اندازه اثر رابطه خودتنظیمی با عملکرد تحصیلی با رویکرد روزنتال و روبین در

۲۷ مطالعه

مطالعه	(r)	Z_{r_i}	n	مقدار p	مطالعه	(r)	Z_{r_i}	n	مقدار p
۱	۰/۵۵	۰/۶۱	۳۴۸	۰/۰۰۰	۱۵	۰/۴۶	۰/۴۹	۶۰	۰/۰۰۰
۲	۰/۲۷	۰/۲۷	۱۲۰	۰/۰۰۳	۱۶	۰/۴۰	۰/۴۲	۸۰	۰/۰۰۰
۳	۰/۱۶	۰/۱۶	۲۵۰	۰/۰۰۱	۱۷	۰/۲۲	۰/۲۲	۱۰۴	۰/۰۰۲
۴	۰/۱۷	۰/۱۷	۴۶	۰/۰۰۲	۱۸	۰/۵۲	۰/۵۷	۲۶۱	۰/۰۰۰
۵	۰/۳۲	۰/۳۳	۱۳۰	۰/۰۰۰	۱۹	۰/۰۸	۰/۰۸	۱۳۰۰	۰/۰۰۴
۶	۰/۴۳	۰/۴۶	۱۵۰	۰/۰۰۰	۲۰	۰/۴۸	۰/۵۲	۴۰۱۰	۰/۰۰۰
۷	۰/۰۴	۰/۰۴	۳۰۰	۰/۰۳۹	۲۱	۰/۱۵	۰/۱۵	۱۴۰۰	۰/۰۱۷
۸	۰/۱۵	۰/۱۵	۱۶۶	۰/۰۰۰	۲۲	۰/۰۲	۰/۰۲	۱۴۰	۰/۰۶۸
۹	۰/۱۸	۰/۱۸	۳۱۵	۰/۰۰۱	۲۳	۰/۴۶	۰/۴۹	۶۰	۰/۰۰۰
۱۰	۰/۴۸	۰/۵۲	۶۰	۰/۰۰۰	۲۴	۰/۵۲	۰/۵۷	۴۰	۰/۰۰۰
۱۱	۰/۰۹	۰/۰۹	۴۳۵	۰/۰۰۶	۲۵	۰/۲۸	۰/۲۸	۴۳۵	۰/۰۰۰
۱۲	۰/۱۶	۰/۱۶	۲۵۰	۰/۰۰۰	۲۶	۰/۵۸	۰/۶۶	۳۲۰	۰/۰۰۰
۱۳	۰/۲۱	۰/۲۱	۳۸۴	۰/۰۰۰	۲۷	۰/۴۱	۰/۴۳	۱۲۰	۰/۰۰۰
۱۴	۰/۷۴	۰/۹۵	۳۰۰	۰/۰۰۰					

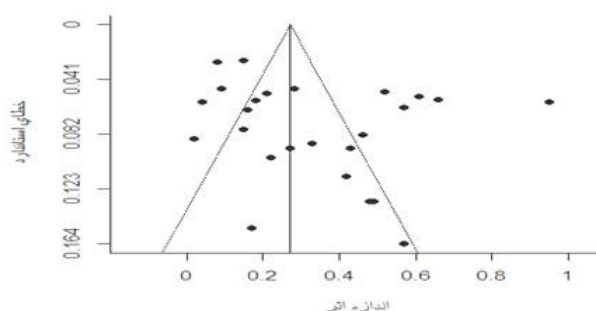
سوگیری دارد. نمودار کیفی تقریباً متقارن شکل ۲ نشان می‌دهد که سوگیری انتشار قابل ملاحظه‌ای رخ نداده است، در نتیجه تعمیم نتایج با دقت بالا صورت می‌گیرد، **اشترن و همکاران (۲۰۰۵)**.

۲.۴. آزمون فرضیه پژوهش. هم‌چنان که گفته شد فرضیه پژوهش حاضر عبارت است از: “بین خودتنظیمی با عملکرد تحصیلی فراگیرندگان رابطه وجود دارد”. چنان‌چه متوسط کلی اندازه اثر را با پارامتر (θ) نشان دهیم، آنگاه آزمون فرضیه پژوهش معادل با آزمون فرضیه‌های فازی زیر است:

فرضیه صفر: متوسط کلی اندازه اثر (θ) “تقریباً” صفر است

فرضیه مقابل: متوسط کلی اندازه اثر (θ) “تقریباً” بزرگ‌تر از صفر است

یعنی



شکل ۲: نمودار کیفی مربوط به سوگیری انتشار برآورد اندازه اثر میزان خودتنظیمی در عملکرد تحصیلی

$$\begin{cases} \tilde{H}_0 : \theta \cong 0 \\ \tilde{H}_1 : \theta \succ 0 \end{cases}$$

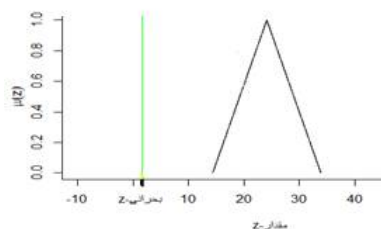
جدول ۳ نتایج به دست آمده از آزمون فرضیه را با استفاده از نرم افزار R نشان می‌دهد.

جدول ۳: اندازه اثر ترکیبی پژوهش‌ها مربوط به اندازه اثر میزان خودتنظیمی در عملکرد تحصیلی با حالت فرضیه فازی

اندازه اثر ترکیبی ($\bar{Z}_{r_i}$)	$\bar{Z}$	$SE(\bar{z}_{r_i})$	درجه اطمینان رد فرضیه صفر (اندازه اثر)	χ^2	df	$p_value_{(\chi^2)}$
۰/۲۷	$(14/31, 24/09, 33/86)_T$	۰/۰۱	$D_{RH_0} = 1$	۴۵/۰۵	۲۶	۰/۰۰۰

جدول ۳ حاکی از آن است که فرضیه صفر پژوهش با درجه اطمینان یک (صد در صد) رد می‌شود. بدین معنی که بر اساس مشاهدات حاصله، و در سطح معنی داری ۵ درصد بین خودتنظیمی و عملکرد تحصیلی با درجه اطمینان صد در صد رابطه معناداری وجود دارد.

همچنین نتایج جدول ۳ حاکی از این است که آزمون نا همگنی (χ^2) دارای مقادیر معنادار است که بیانگر ناهمگنی تحقیقات مورد بررسی می‌باشد و ناهمگن بودن مطالعات نشان می‌دهد که متغیرهای مداخله گر بین این دو متغیر خودتنظیمی و عملکرد تحصیلی وجود دارد. جهت ارزیابی عوامل ایجاد کننده ناهمگونی از مدل فرا رگرسیون فازی با حالت داده دقیق و فرضیه



شکل ۳: نمودار تابع عضویت آماره آزمون $\tilde{Z}$ مربوط به اندازه اثر میزان خودتنظیمی یادگیری در عملکرد تحصیلی با فرضیه فازی

فازی استفاده گردید. برای دست یابی به این مهم، سال انجام مطالعات و اندازه نمونه به عنوان عوامل محتمل ایجاد کننده ناهمگونی از عوامل مختلف در مدل وارد شدند. جدول ۴ نتایج حاصل از برازش مدل رگرسیون خطی فازی متقارن چندگانه اندازه اثر مطالعات روی متغیرهای حجم نمونه و سال انتشار را بر اساس h های مختلف نشان می‌دهد.

جدول ۴: نتایج فرا رگرسیون فازی در ایجاد ناهمگونی رابطه بین خودتنظیمی و عملکرد تحصیلی

h	a_0	a_1	a_2	s_0	s_1	s_2	ابهام کل مدل (z)	MDM	IC
۰/۲	۱۱/۱۶	۰/۰۰۰۷	-۰/۰۰۷۸	۰/۳۳	۰/۰۰۷	۰/۰۰	۳۰/۷۹	۰/۸۶	۰/۸۸
۰/۴	۱۱/۱۶	۰/۰۰۰۷	-۰/۰۰۷۸	۰/۴۴	۰/۱۰	۰/۰۰	۴۱/۰۶	۰/۸۸	۰/۸۹
۰/۵	۱۱/۱۶	۰/۰۰۰۷	-۰/۰۰۷۸	۰/۵۳	۰/۱۲	۰/۰۰	۴۹/۲۷	۰/۸۸	۰/۸۹
۰/۶	۱۱/۱۶	۰/۰۰۰۷	-۰/۰۰۷۸	۰/۶۷	۰/۱۵	۰/۰۰	۶۱/۵۹	۰/۸۹	۰/۸۹
۰/۸	۱۱/۱۶	۰/۰۰۰۷	-۰/۰۰۷۸	۱/۳۴	۰/۳۱	۰/۰۰	۱۲۳/۱۹	۰/۸۹	۰/۸۹

چنان که از نتایج جدول ۴ ملاحظه می‌شود، با بزرگ شدن مقدار h سطح اعتبار مدل بالا می‌رود و در عین حال این امر باعث افزایش ابهام کل مدل نیز می‌شود. لذا برای انتخاب h صحیح به میزان ابهام کل مدل نیز باید توجه داشت. در داده‌های مورد بررسی سطح اعتبار $h = ۰/۵$ به سبب شاخص IC و MDM بالا و افزایش اندک ابهام مدل معقول به نظر می‌رسد. بنا بر این

در حالت متقارن مدل رگرسیونی به صورت زیر خواهد بود

$$y = \text{سال نشر } (-0.007, 0.00) + \text{حجم نمونه } (0.0007, 0.12) + (11.16, 0.53)$$

معادله بالا نشان می‌دهد که با ثابت نگاه داشتن سایر متغیرها با افزایش یک واحد به متغیر حجم نمونه اندازه اثر حدوداً به اندازه ۰/۰۰۰۷ افزایش، و با افزایش یک واحد به متغیر سال نشر، دقیقاً به مقدار ۰/۰۰۷ - اندازه اثر کاهش می‌یابد.

۵. بحث و نتیجه‌گیری

نتایج حاصل از پژوهش حاضر نشان داد که میانگین اندازه اثر (اثرات ترکیب ثابت) رابطه متغیر خودتنظیمی با عملکرد تحصیلی در نمونه مورد پژوهش معادل ۰/۲۷ می‌باشد، از آنجا که اندازه اثر برآورد شده در سطح ۵ درصد با درجه پذیرش ۱۰۰ درصد معنادار می‌باشد، لذا متغیر خودتنظیمی به تنهایی حدوداً به اندازه ۰/۰۷ از تغییرات عملکرد تحصیلی فراگیرندگان را تبیین می‌کند. بر اساس مدل کوهن، کوهن (۱۹۷۷)، این رابطه در حد متوسط ارزیابی می‌شود. هم چنین نتایج فرا رگرسیون فازی پژوهش نشان داد متغیر اندازه نمونه، سال انتشار به عنوان متغیر تعدیل کننده بین متغیر خودتنظیمی و عملکرد تحصیلی نقش ایفا می‌کنند. نتایج حاصل از بررسی رابطه خودتنظیمی و عملکرد تحصیلی نشان می‌دهد که نتایج مطالعه حاضر با مطالعات پژوهش‌های شماره ۷، ۲۱ و ۲۲ در جدول ۲ ناهمسو و با سایر مطالعات ارائه شده در جدول ۲ همسو هستند، جو (۱۹۹۷).

مراجع

اسدی، م. (۱۳۹۴). طراحی و اعتباریابی الگوی والدگری موفق در نمونه هایی از خانواده‌های ایرانی: پژوهشی ترکیبی مبتنی بر ریاضیات فازی. رساله دکتری مشاوره، دانشگاه خوارزمی، تهران.

- خدایی، ا. (۱۳۸۸). رگرسیون خطی فازی و کاربرد آن در پژوهش‌های اجتماعی، *مجله مطالعات اجتماعی/ایران*، ۳، ۸۲-۹۸.
- طاهری، س.م.، ماشین چی، م. (۱۳۸۷). *مقدمه‌ای بر احتمال و آمار فازی*، انتشارات دانشگاه شهید باهنر کرمان.
- هومن، ح.ع. (۱۳۸۷). *راهنمای عملی در پژوهش فراتحلیل*، انتشارات سمت.
- Borenstein, M., Hedges, L., Higgins, J., & Rothstein, H.R. (2009), *Computing Effect Sizes for Introduction to Meta-Analysis*, Academic Press, New York.
- Cohen J. (1977), *Statistical Power Analysis for the Behavioral Sciences*, Academic Press, New York.
- Field A.P. (2001), Meta-Analysis of Correlation Coefficients: A Monte Carlo Comparison of Fixed- and Random-Effects Methods, *Psychological Methods*. 6, 161 - 180.
- Joe H. (1997), *Multivariate Models and Dependence Concepts*, Chapman & Hall, London.
- Pigott T.D. (2012), *Statistics for Social and Behavioral Sciences*, Springer, New York.
- Pourahmad S., Ayatollahi, S.M.T. & Taheri S.M. (2011), Fuzzy logistic regression: A new possibilistic model, and its application in clinical vague status, *Iranian Journal of Fuzzy Systems*. 8, 1-17.
- Stanly, T.D. & Jarrel, S.B. (1989). Meta regression analysis: A quantitative method of literature surveys, *Journal of Economic Surveys*. 19, 299 - 308.

- Sterne J.A.C., Becker B.J., & Egger M. (2005), The funnel plot, In: *Publication Bias in Meta-Analysis: Prevention, Assessment and Adjustment*, Academic Press, New York, pp: 75-98.
- Taheri S.M. & Arefi M. (2009), Testing fuzzy hypotheses based on fuzzy test statistics, *Soft Computing*. 13, 617 - 625.



آزمون فرضیه در محیط نادقیق با استفاده از روش بیز امکانی

ملیحه عابدینی^{* 1}

mabedini1073@gmail.com

چکیده. آزمون فرضیه‌های آماری از اهمیت زیادی برای تصمیم‌گیری در مسائل کاربردی برخوردار است. در روش‌های مرسوم آزمون فرضیه‌های آماری، داده‌ها، فرضیه‌ها، پارامترها و دیگر اجزای موجود در مسئله دقیق در نظر گرفته می‌شوند. در حالی که در علوم کاربردی ممکن است در میزان و یا تعریف یک یا چند جزء مسئله ابهام وجود داشته باشد. در این حالت روش‌های معمولی کارایی ندارند و نیازمند تعمیم می‌باشند. در این مقاله، یک نگرش بیزی برای آزمون فرضیه آماری را معرفی و بررسی می‌کنیم. در این نگرش، بر پایه یک مدل آماری و یک توزیع پیشین امکانی، یک توزیع پسین امکانی محاسبه می‌شود و آزمون فرضیه به صورت مناسب تعریف می‌شود. این ایده براساس تعمیمی از رویکرد لاپوینت و بویی (۲۰۰۰) برای الگوی بیز امکانی استوار است.

۱. مقدمه

استنباط بیزی برای مدل‌های آماری پارامتری بر پایه دو اصل زیر استوار است:
۱- متغیر X با مدل آماری $f(x|\theta)$ و پارامتر θ با توزیع پیشین $\pi(\theta)$ ماهیتی تصادفی دارند.
۲- داده‌های مرتبط با متغیر تصادفی X دقیق گزارش می‌شوند.
اصول فوق، اصول مهمی در نگرش بیز کلاسیک می‌باشند. اما در بسیاری از بررسی‌های آماری ممکن است مشاهدات، فرضیه‌های مورد آزمون یا پارامترهای مورد بررسی دقیق نباشند. در این موارد، روش‌های آزمون فرضیه آماری در حالت کلاسیک کارایی و اعتبار لازم را ندارند. نظریه مجموعه‌های فازی شیوه‌های مناسبی را برای صورت‌بندی و تحلیل مسائل آماری در این حالت‌های نادقیق پیشنهاد می‌کند.

در این راستا، در این مقاله، یک نگرش بیزی برای آزمون فرضیه آماری را معرفی و بررسی می‌کنیم. در این نگرش، بر پایه یک مدل امکانی و یک توزیع پیشین امکانی، یک توزیع پسین امکانی محاسبه می‌شود و آزمون فرضیه به صورت مناسب تعریف می‌شود. این ایده، براساس رویکرد لاپوینت و بویی [۶] برای الگوی بیز امکانی است. در واقع، لاپوینت و بویی براساس یک مدل امکانی برای متغیر مورد بحث و یک توزیع پیشین امکانی برای پارامتر مورد نظر، یک

واژگان کلیدی. آزمون بیزی، توزیع پسین امکانی، توزیع پیشین امکانی، نگرش بیز امکانی.
* سخنران

توزیع پسین امکانی تحت داده‌های معمولی تعریف نموده‌اند. در مقاله حاضر آزمون فرضیه آماری براساس تابع پسین امکانی تعمیم داده می‌شود. استنباط بیزی در محیط فازی توسط برخی نویسندگان بررسی شده است. در این جا، به‌طور مختصر، به تعدادی از این تحقیقات اشاره می‌کنیم. برآورد بیزی برای حالتی که داده‌های موجود فازی باشند، توسط گیل و همکاران [۵] تعمیم یافته است. شناتر [۷] به بررسی استنباط بیزی در حالت‌های مختلفی که داده‌ها دقیق یا فازی و پارامتر مورد نظر نیز دقیق یا فازی باشند، پرداخته است. یک دیدگاه متفاوت در بحث استنباط بیز امکانی توسط لاپوینت و بوبی [۶] تحت تابع پسین امکانی مورد بررسی قرار گرفته است. تعمیمی از روش لاپوینت و بوبی زمانی که داده‌ها فازی باشند توسط عارفی و طاهری [۲] مورد مطالعه قرار گرفته است. یک دیدگاه بیزی براساس تابع چگالی احتمالی و توزیع پیشین امکانی تحت داده‌های فازی توسط عارفی و طاهری [۱] مورد بررسی و مطالعه قرار گرفته است. همچنین طاهری و بهبودیان [۸] آزمون فرضیه‌های فازی تحت نگرش بیزی را در دو حالتی که داده‌ها دقیق یا فازی باشند، تعمیم داده‌اند. این مقاله شامل ۵ بخش است. در بخش ۲، مفاهیم ضروری مورد استفاده در مقاله مرور می‌گردند. سپس در بخش ۳، فرضیه‌های امکانی معرفی می‌شوند. در بخش ۴، به محاسبه یک روش بیز امکانی برای آزمون فرضیه پرداخته می‌شود. در بخش ۵، یک مثال عددی ارائه می‌شود.

۲. مفاهیم اولیه

۱.۲. توزیع امکان. فرض کنید Ω یک مجموعه مرجع و B یک میدان سیگمایی از زیرمجموعه‌های Ω باشد.

تعریف ۱.۲. تابع مجموعه‌ای $\Pi : B \rightarrow [0, 1]$ را یک اندازه/امکان بر (Ω, B) گوئیم، هرگاه در شرایط زیر صدق کند:

1. $\Pi(A) \geq 0$
2. $\Pi(\Omega) = 1$
3. $\Pi(\bigcup_{i=1}^{\infty} A_i) = \sup_i \Pi(A_i), \forall A_i \in B$

تعریف ۲.۲. فرض کنید Π یک اندازه/امکان بر (Ω, B) باشد. در این صورت تابع عضویت $\pi(x)$ را که متناظر با یک مجموعه فازی از Ω است و به صورت زیر تعریف می‌شود، تابع توزیع امکان متناظر با اندازه/امکان Π می‌نامیم.

$$\Pi(A) = \sup\{\pi(x) | x \in A\}$$

تعریف ۳.۲. توزیع امکان توأم متغیر X (که در U مقدار می‌گیرد) و Y (که در V مقدار می‌گیرد) به وسیله $\pi_{(X,Y)}(u,v)$ مشخص می‌شود و امکان این که به‌طور همزمان $X = u$ و $Y = v$ باشد را نشان می‌دهد. توزیع امکان شرطی Y به شرط $X = u$ به وسیله $\pi_{Y|X}(v|u)$ مشخص می‌شود و امکان این پیشامد که $Y = v$ وقتی می‌دانیم $X = u$ است را نشان می‌دهد.

۲.۲. رابطه‌های فازی و ترکیب آن‌ها.

تعریف ۴.۲. یک رابطه فازی بین مجموعه‌های مرجع U و V ، یک مجموعه فازی از فضای $U \times V$ است.

تعریف ۵.۲. فرض کنید $Q \in L(U \times V)$ و $R \in L(V \times W)$ دو رابطه فازی باشند. ترکیب Q و R رابطه فازی $P = Q \circ R \in L(U \times W)$ با تابع عضویت زیر را نتیجه می‌دهد:

$$\mu_P(u, w) = S_V[T(Q(u, v), R(v, w))],$$

که در آن T هر T -نرم و S هر T -هم‌نرم دلخواه است.

توجه ۶.۲. در این مقاله همواره از عملگر $\sup$ برای T -هم‌نرم استفاده می‌شود. همچنین فرض می‌شود T یک T -نرم پیوسته باشد، که در این مقاله T -نرم حاصل ضرب را در نظر گرفته‌ایم.

تعریف ۷.۲. فرض کنید $R \in L(U \times V)$ یک رابطه فازی و $A \in L(U)$ یک مجموعه فازی باشد. ترکیب R و A مجموعه فازی $B = A \circ R \in L(V)$ با تابع عضویت زیر را نتیجه می‌دهد:

$$\mu_B(v) = S_U[T(A(u), R(u, v))]. \quad (۱.۲)$$

تعریف ۸.۲. فرض کنید متغیر X دارای توزیع امکان $\pi(x|\theta)$ و پارامتر θ دارای توزیع پیشین امکانی $\pi(\theta)$ باشند. آنگاه توزیع امکان توأم متغیرهای X و θ به صورت زیر تعریف می‌شود:

$$\pi(x, \theta) = T(\pi(\theta), \pi(x|\theta)). \quad (۲.۲)$$

توجه ۹.۲. بر پایه معادله ۱.۲، توزیع حاشیه‌ای امکانی X به صورت زیر بدست می‌آید:

$$\pi(x) = \sup_{\Theta} [T(\pi(\theta), \pi(x|\theta))] = \sup_{\Theta} [\pi(x, \theta)], \quad (۳.۲)$$

که در آن Θ فضای پارامتر است.

تعریف ۱۰.۲. مجموعه $X = (X_1, \dots, X_n)$ را با امکان $\pi(X|\theta)$ در نظر بگیرید. فرض کنید پارامتر θ دارای تابع امکان پیشین $\pi(\theta)$ باشد. توزیع امکان پسین براساس عملگر $\psi_T(\cdot, \cdot)$ به صورت زیر تعریف می‌شود:

$$\pi(\theta|X) = \psi_T(\pi(X), \pi(X, \theta)),$$

که در آن

$$\pi(X, \theta) = T(\pi(\theta), \pi(X|\theta)),$$

$$\pi(X) = \sup_{\theta \in \Theta} (\pi(X, \theta)),$$

که $\psi_T(\cdot, \cdot)$ عملگر R -استانزام می‌باشد، که براساس لاپوینت و بویی [۶] به صورت زیر محاسبه می‌شود:

$$\psi_T(a, c) = \sup\{x \in [0, 1] : T(a, x) \leq c\}.$$

توجه ۱۱.۲. عملگر ψ_T برای T -نرم‌های مختلف به صورت زیر محاسبه می‌شود:

$$\begin{aligned} \psi_{T_1}(a, c) &= \begin{cases} 1 & a \leq c, \\ c & a > c, \end{cases} & T_1(a, b) &= \min(a, b), \\ \psi_{T_2}(a, c) &= \begin{cases} 1 & a \leq c, \\ 1 + c - a & a > c, \end{cases} & T_2(a, b) &= \max(0, a + b - 1), \\ \psi_{T_3}(a, c) &= \begin{cases} 1 & a \leq c, \\ \frac{c}{a} & a > c, \end{cases} & T_3(a, b) &= a.b, \\ \psi_{T_4}(a, c) &= \begin{cases} 1 & a \leq c, \\ \frac{2c-a.c}{a+c-a.c} & a > c, \end{cases} & T_4(a, b) &= \frac{a.b}{1 + (1-a).(1-b)}, \\ \psi_{T_5}(a, c) &= \begin{cases} 1 & a \leq c, \\ \frac{a.c}{a-c+ac} & a > c, \end{cases} & T_5(a, b) &= \frac{a.b}{a + b - a.b}. \end{aligned}$$

۳. فرضیه‌های امکانی

یکی از مسائل مهم در آزمون‌های آماری، تنظیم و صورت‌بندی فرضیه‌هاست. گاهی، ماهیت فرضیه‌ها به گونه‌ای است که به شکل یک فرضیه دقیق فرمول‌بندی نمی‌شوند. برگر و دلامپادی [۴] غیر واقعی بودن فرضیه‌های دقیق را این گونه توضیح می‌دهند:

”به ندرت اتفاق می‌افتد و یا شاید غیر ممکن است که در عالم واقعی، فرضیه‌ای به شکل $\theta = \theta_0$ مدل‌بندی شود.“ مثلاً، فرضیه‌هایی به شکل زیر در آمار خیلی رایج هستند:

H_0 : ویتامین C اثری بر سرماخوردگی ندارد.

واضح است که این فرضیه را نمی‌توان به شکل یک فرضیه دقیق در نظر گرفت.

تعریف ۱.۳. در یک مسأله آزمون فرضیه، هر فرضیه به شکل

$H: \theta$ به صورت $H(\theta)$ است.

را یک فرضیه امکانی گوئیم، که در آن $H(\theta)$ یک تابع امکان روی فضای پارامتر Θ است.

توجه ۲.۳. در حالت معمولی، فرض $H: \theta \in \Theta$ را می‌توان یک فرضیه با تابع امکان زیر در نظر گرفت:

$$H(\theta) = \begin{cases} 1 & \theta \in \Theta, \\ 0 & \theta \notin \Theta. \end{cases}$$

مثال ۳.۳. فرض کنید θ پارامتر نسبت در یک توزیع دوجمله‌ای باشد. فرض ” θ تقریباً ۰.۵ است“، یک فرضیه امکانی است که می‌تواند به صورت زیر تعریف شود:

$$H(\theta) = \begin{cases} 2\theta & 0 \leq \theta < \frac{1}{2}, \\ 2 - 2\theta & \frac{1}{2} \leq \theta < 1. \end{cases}$$

۴. آزمون فرضیه امکانی

فرض کنید X یک متغیر تصادفی باشد که دارای تابع امکان $\pi(x|\theta), \theta \in \Theta$ است. فرض کنید دو تابع عضویت $H_0(\theta)$ و $H_1(\theta)$ داده شده‌اند و می‌خواهیم بر پایه روش بیز امکانی، فرضیه‌های امکانی زیر را آزمون کنیم:

$$\begin{cases} \theta \text{ به صورت } H_0(\theta) \text{ است:} \\ \theta \text{ به صورت } H_1(\theta) \text{ است:} \end{cases}$$

در ابتدا، به طور خلاصه، بعضی از تعاریف مورد نیاز از نظریه تصمیم را بیان می‌کنیم. فرض کنید $\mathcal{A} = \{a_0, a_1\}$ فضای عمل می‌باشد که در آن a_0 نشان‌دهنده پذیرش فرض H_0 و a_1 نشان دهنده پذیرش H_1 است.

تابع امکان پیشین $\pi(\theta)$ را برای θ در نظر بگیرید و فرض کنید $\pi(x|\theta)$ تابع امکان شرطی X به شرط مقدار ثابت $\theta \in \Theta$ باشد. امکان توأم θ به شرط $X = x$ ، امکان پسین θ نامیده می‌شود و با $\pi(\theta|x)$ نشان داده می‌شود.

۱.۴. آزمون بیز بدون تابع زیان. در این قسمت، مجموعه‌ای از اعمال با دو نقطه a_0 و a_1 را در نظر می‌گیریم، که a_0 مطابق با پذیرش H_0 و a_1 متناظر با رد H_1 است. در این جا با فرضیه‌های امکانی مواجه هستیم.

تعریف ۱.۴. در آزمون بیز، فرضیه H_0 رد می‌شود، اگر و تنها اگر، تابع امکان پسین تحت H_0 ، کمتر باشد از تابع امکان پسین تحت H_1 ، یعنی

$$\sup[T(\pi(\theta|x), H_0(\theta))] < \sup[T(\pi(\theta|x), H_1(\theta))]$$

در بعضی از شرایط، ممکن است دو مقدار α_0 و α_1 ، محاسبه شده در ۱.۴ و ۲.۴ به هم نزدیک باشند، بنابراین برای تصمیم‌گیری در چنین شرایطی از تعریف ۲.۴ استفاده می‌کنیم.

$$\alpha_0 = \sup[T(\pi(\theta|x), H_0(\theta))] \quad (1.4)$$

$$\alpha_1 = \sup[T(\pi(\theta|x), H_1(\theta))] \quad (2.4)$$

تعریف ۲.۴. فرض کنید در تعریف ۱.۴، آزمون بیز فرضیه H_0 را قبول کند، آنگاه مقدار $\frac{\alpha_0}{\alpha_0 + \alpha_1}$ ، درجه پذیرش H_0 در مقابل H_1 گفته می‌شود.

۵. مثال عددی

در این جا برای توضیح موارد گفته شده مثالی را بیان می‌کنیم. این مثال از منبع لاپوینت و بویی [۶] انتخاب شده است. در این جا داده‌ها به صورت معمولی می‌باشند اما مدلی که به داده‌ها برازش می‌دهیم یک مدل امکان می‌باشد.

مثال ۱.۵. اولین قدم، تعیین مدل توزیع امکان پیشین $\pi(x)$ است. فرض کنید که بنا بر گزارش یک خبره، بدانیم مقادیر متغیر X در بازه $[5, 8]$ تغییر می‌کند. علاوه بر این، در بعضی حالت‌های خاص مقدار متغیر کمتر از ۵ است اما هرگز کمتر از ۲ نمی‌باشد و بیشتر از ۸ است اما هرگز بیشتر

از 10 نیست. با استفاده از عدد فازی دوزنقه‌ای توزیع امکان پیشین را به صورت زیر مدل بندی می‌کنیم. (شکل ۱)

$$\pi(\theta) = \begin{cases} \frac{\theta-2}{3} & 2 < \theta < 5 \\ 1 & 5 \leq \theta \leq 8 \\ \frac{10-\theta}{2} & 8 < \theta < 10 \end{cases}$$

از یک عدد فازی مثلثی به صورت زیر برای مدل بندی $\pi(x|\theta)$ استفاده می‌کنیم که پهنای تکیه‌گاه $\pi(x|\theta)$ با افزایش θ ، افزایش پیدا می‌کند. (شکل ۲)

$$\pi(x|\theta) = \begin{cases} \frac{3x-6}{\theta} - 1 & \frac{\theta}{3} + 2 < x \leq \frac{2\theta}{3} + 2 \\ \frac{6-3x}{\theta} + 3 & \frac{2\theta}{3} + 2 < x \leq \theta + 2 \end{cases} \quad (۱.۵)$$

در شکل ۴، $\pi(x|\theta)$ به عنوان تابعی از θ به ازای $\theta = 4, 6, 8$ رسم شده است و در شکل ۳، $\pi(\theta|x)$ به عنوان تابعی از θ برای $X = 4, 6, 8$ رسم شده است.

اینک حالتی را در نظر بگیرید که پیش‌بینی نسبتاً کوچک باشد یعنی $X = 4$ ، با استفاده از رابطه ۱.۵ خواهیم داشت:

$$\pi(4|\theta) = \begin{cases} -\frac{6}{\theta} + 3 & 2 < \theta < 3 \\ \frac{6}{\theta} - 1 & 3 \leq \theta < 6 \end{cases}$$

با استفاده از قاعده، $T(a, b) = a.b$ رابطه زیر نتیجه می‌شود:

$$\pi(\theta|4) = \begin{cases} -\frac{3(\theta^2-4\theta+4)}{\theta} & 2 < \theta < 3 \\ \frac{-\theta^2+8\theta-12}{\theta} & 3 \leq \theta < 5 \\ \frac{18}{\theta} - 3 & 5 \leq \theta < 6 \end{cases}$$

اکنون می‌خواهیم آزمون

$$H_0: \theta \simeq 2$$

$$H_1: \theta \simeq 4$$

که $\theta \simeq 2$ به معنی ” θ تقریباً برابر ۲ است.“ و $\theta \simeq 4$ به معنی ” θ تقریباً برابر ۴ است.“ را انجام دهیم. توابع عضویت به صورت زیر هستند:

$$H_0(\theta) = \begin{cases} \theta - 1 & 1 < \theta \leq 2 \\ 3 - \theta & 2 < \theta < 3 \end{cases}$$

$$H_1(\theta) = \begin{cases} \theta - 3 & 3 < \theta \leq 4 \\ \frac{7-\theta}{3} & 4 < \theta < 7 \end{cases}$$

بنابراین با توجه به فرض صفر خواهیم داشت:

$$T(H_0(\theta), \pi(\theta|4)) = \frac{3(\theta^2 - 4\theta + 4)(3 - \theta)}{\theta}$$

$$\sup_{0 < \theta < 1} T(H_0(\theta), \pi(\theta|4)) = \frac{3(\theta^2 - 4\theta + 4)(3 - \theta)}{\theta} = 0.166 \quad (۲.۵)$$

و همچنین براساس فرض مقابل نتیجه می‌شود:

$$T(H_1(\theta), \pi(\theta|4)) = \begin{cases} \frac{(\theta-3)(-\theta^2+8\theta-12)}{\theta} & 3 < \theta < 4 \\ \frac{(\frac{7-\theta}{3})(-\theta^2+8\theta-12)}{\theta} & 4 < \theta < 5 \\ \frac{7-\theta}{3} (\frac{18}{\theta} - 3) & 5 < \theta < 6 \end{cases}$$

$$\sup_{0 < \theta < 1} T(H_1(\theta), \pi(\theta|4)) = 1 \quad (3.5)$$

بنابراین با توجه به ۲.۵ و ۳.۵ نتیجه می‌شود:

$$\sup_{0 < \theta < 1} T(H_0(\theta), \pi(\theta|4)) < \sup_{0 < \theta < 1} T(H_1(\theta), \pi(\theta|4))$$

لذا فرض H_0 رد می‌شود.

اگر فرضیه‌های صفر و یک را به صورت زیر در نظر بگیریم:

$$H_0 : \theta \simeq 6$$

$$H_1 : \theta \simeq 9$$

که $\theta \simeq 6$ به معنی " θ تقریباً برابر ۶ است." و $\theta \simeq 9$ به معنی " θ تقریباً برابر ۹ است." را انجام دهیم. توابع عضویت به صورت زیر هستند:

$$H_0(\theta) = \begin{cases} \frac{\theta-4}{\frac{9-2}{3}} & 4 < \theta \leq 6 \\ \frac{9-2}{3} & 6 < \theta < 9 \end{cases}$$

$$H_1(\theta) = \begin{cases} \frac{\theta-7}{\frac{16-2}{7}} & 7 < \theta \leq 9 \\ \frac{16-2}{7} & 9 < \theta < 16 \end{cases}$$

بنابراین با توجه به فرض صفر خواهیم داشت:

$$T(H_0(\theta), \pi(\theta|4)) = \frac{(-\theta^2 + 8\theta - 12)(\theta - 4)}{2\theta}$$

$$\sup_{0 < \theta < 1} T(H_0(\theta), \pi(\theta|4)) = 0.3 \quad (4.5)$$

و همچنین براساس فرض مقابل نتیجه می‌شود:

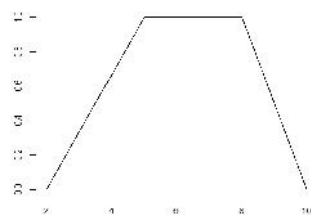
$$T(H_1(\theta), \pi(\theta|4)) = (\frac{18}{\theta} - 3)(\frac{\theta - 4}{2})$$

$$\sup_{0 < \theta < 1} T(H_1(\theta), \pi(\theta|4)) = 0 \quad (5.5)$$

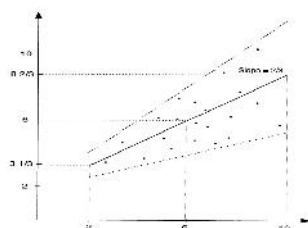
بنابراین با توجه به ۴.۵ و ۵.۵ نتیجه می‌شود:

$$\sup_{0 < \theta < 1} T(H_0(\theta), \pi(\theta|4)) > \sup_{0 < \theta < 1} T(H_1(\theta), \pi(\theta|4))$$

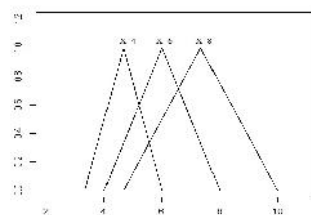
لذا فرض H_0 رد نمی‌شود.



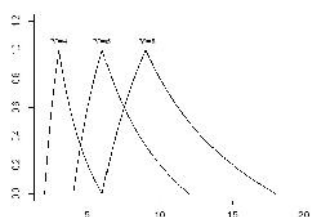
شکل ۱: توزیع امکان پیشین در مثال ۱.۵.



شکل ۲: پیش‌بینی θ در مقابل مقادیر مشاهد شده X در مثال ۱.۵.



شکل ۳: توزیع امکان شرطی $\pi(\theta|x)$ برای $X = 4, 6, 8$ در مثال ۱.۵.



شکل ۴: توزیع امکان شرطی $\pi(\theta|x)$ برای $\theta = 4, 6, 8$ در مثال ۱.۵.

۶. نتیجه گیری

در این مقاله، شیوه جدیدی برای آزمون فرضیه در حالت‌های نادقیق، معرفی و بررسی گردید. اساس روش پیشنهاد شده مبتنی بر الگوی بیز امکانی ساخته شده برای پارامتر و داده‌های مورد نظر است. روش کار بدین صورت است که تعریف آزمون بیزی با استفاده از داده‌ها و فرضیه‌های معمولی را به حالتی که داده‌ها و فرضیه‌ها با استفاده از تابع امکان بیان شده باشند، تعمیم می‌دهیم. برای این منظور از T -نرم حاصل ضرب و S -نرم $\sup$ استفاده می‌کنیم.

مراجع

۱. م. عارفی و س. م. طاهری، برآورد بیزی براساس توزیع پسین امکانی با داده‌های فازی، مجله اندیشه آماری، شماره ۱، ۷۷-۸۵.
2. Arefi, M. and Taheri, S.M. (2016), Possibilistic Bayesian inference based on fuzzy data, *International Journal of Machine Learning and Cybernetics*, 7, 753-763.
3. Berger, J.O. (1985), *Statistical Decision Theory and Bayesian Analysis*, Springer, New York.
4. Berger, O. and Delampady, M. (1987), Testing Precise Hypotheses, *Statistical Science*, 2, 317-352.
5. Gil, M., Corral, N. and Gil, P. (1985), The fuzzy decision problem: An approach to the point estimation problem with fuzzy information, *European Journal of Operational Research*, 22, 26 - 34.
6. Lapointe, S. and Bobee, B. (2000), Revision of Possibility distributions: A Bayesian inference pattern, *Fuzzy Sets and Systems*, 116, 119-140.
7. Schnatter, S. (1993), On fuzzy Bayesian inference, *Fuzzy Sets and Systems*, 60, 41-58.
8. Taheri, S.M. and Behboodian, J. (2006), On Bayesian approach to fuzzy testing hypothesis with fuzzy data, *Italian Journal of Pure and Applied Mathematics*, 19, 139-154.
9. Viertl, R. (2008), Foundations of fuzzy Bayesian inference, *Journal of Uncertain Systems*, 2, 187-191.



استوارسازی تابع عضویت فازی فاصله مبنا برای ماشین بردار پشتیبان

ماندانا محمدی* و مجید سرمد

گروه آمار، دانشگاه فردوسی مشهد
manda.mohammadi@stu.um.ac.ir
sarmad@um.ac.ir

چکیده. ماشین بردار پشتیبان یکی از الگوریتم های موفق در زمینه طبقه بندی است که به دلیل داشتن ویژگی های منحصر به فرد، استفاده از آن گسترش چشم گیری داشته است. در عین حال کارایی این روش تحت تأثیر مشاهدات دور افتاده قرار گرفته و دچار تنزل می شود. در این مقاله با استفاده از مفهوم فازی، به وسیله تعریف یک "تابع عضویت" به مقابله با تأثیر سوء وجود این گونه مشاهدات پرداخته شده است. ساختار تابع عضویت بکار رفته بر مبنای معیار "فاصله ی" نقاط تا مرکز کلاس مربوطه بنا نهاده شده است. غالباً مرکز هر کلاس با استفاده از میانگین حسابی محاسبه می شود که متأسفانه، تأثیر پذیری آن نسبت به وجود مشاهدات دور افتاده به اثبات رسیده است. بنابراین، با تکیه بر روش های آمار استوار، مانند برآوردگر مینیمم درمیان کوواریانس^۱ به استوارسازی تابع عضویت فازی می پردازیم. به کارگیری همزمان دو روش "فازی سازی" و "استوارسازی" باعث افزایش دقت و قدرت تعمیم پذیری ماشین بردار پشتیبان می شود که این امر با استفاده از داده های شبیه سازی شده و واقعی مورد آزمایش و تأیید قرار گرفت.

۱. پیش گفتار

ماشین بردار پشتیبان^۱ به عنوان یکی از روش های یادگیری با ناظر است که از تئوری یادگیری آماری نشأت گرفته است [۱]. قابلیت تعمیم پذیری بالا، دلیلی بر به کارگیری گسترده آن در بسیاری از روش های داده کاوی، مهندسی و بیوانفورماتیک شده است [۳]. از دلایل عمده محبوبیت SVM می توان عواملی همچون دقت طبقه بندی، تفسیر آسان از لحاظ هندسی و تئوری زیربنایی قوی نام برد [۱۲]. اما با وجود تمامی خصوصیات ذکر شده، این الگوریتم تحت تأثیر مشاهدات دور افتاده قرار می گیرد که باعث افت عملکرد و کاهش قدرت تعمیم پذیری آن می شود. بر مبنای تعریف هاکینز، مشاهده دور افتاده، داده ای است که نسبت به سایر مشاهدات بصورت قابل ملاحظه ای متفاوت باشد [۲]. برای فائق آمدن بر مشکلات ناشی از حضور مشاهدات دور افتاده،

واژگان کلیدی. ماشین بردار پشتیبان، مشاهدات دور افتاده، فازی سازی، استوارسازی، برآوردگر مینیمم درمیان کوواریانس.
* سخنران

^۱Support vector Machine (SVM)

”فازی سازی“^۲ SVM پیشنهاد شده است. فازی سازی در این مقاله با استفاده از روش ”تابع عضویت“^۳ انجام شده است. این تابع یکی از مفاهیم مهم منطق فازی است که در واقع تعمیم یافته تابع مشخصه، در مجموعه های کلاسیک است. فازی سازی SVM اولین بار توسط لین و ونگ معرفی شد [۵]، که با اختصاص دادن ضرایب وزنی مختلف به هر یک از نمونه های آموزشی با استفاده از مقادیر تابع عضویت، باعث کنترل تأثیر مشاهدات دورافتاده در مکان قرارگیری ابر صفحه جداکننده می شود. اما ذکر این نکته ضروری است که انتخاب تابع عضویت مناسب یک مرحله بسیار مهم در SVM فازی است.

یکی از معیارهای معمول برای شناسایی مشاهدات دورافتاده، محاسبه میزان فاصله آن ها تا مرکز مشاهدات است [۵]. فاصله ”اقلیدسی“^۴ مهم ترین فاصله برای این منظور است که در تابع عضویت معرفی شده توسط لین [۵]، به کار برده شده است. گفتنی است که میانگین حسابی به کار رفته در فاصله اقلیدسی را می توان به عنوان اولین انتخاب برای مرکز هر کلاس در نظر گرفت. اما باید این نکته را در نظر داشت که روش های برآورد آمار کلاسیک مانند برآوردگر میانگین و کوواریانس تحت تأثیر مشاهدات دورافتاده قرار گرفته و کارایی خود را از دست می دهند. بنابراین، برای رهایی از این معضل به کارگیری روش های برآورد استوار که نسبت به نقص فرضیات زیربنایی آمار کلاسیک و حضور مشاهدات دورافتاده حساس نیستند، مفید خواهد بود. در این راستا، می توان از انواع مختلف برآوردهای استوار به عنوان جایگزین مناسبی برای میانگین استفاده کرد ([۷] و [۸]). در این مقاله از برآوردگر ”مینیم دترمینان کوواریانس“^۵ [۹] که به طور خلاصه در بخش ۴.۲ معرفی شده است، استفاده کرده ایم.

به علاوه، برای بررسی امکان وجود روشی با دقت بالاتر نسبت به فاصله اقلیدسی از فاصله ”ماهالانوبیس“^۶ به عنوان جایگزینی برای آن استفاده کردیم. این فاصله، علاوه بر این که فاصله هر نقطه تا مرکز مشاهدات را محاسبه می کند، به طور همزمان کوواریانس داده ها را در محاسبات لحاظ می کند.

در این مقاله، ترکیب همزمان دو مقوله فازی سازی و استوارسازی با استفاده از توابع عضویت طراحی شده بر مبنای برآوردگر استوار مکان و مقیاس MCD، به طور چشمگیری باعث کاهش تأثیر مشاهدات دورافتاده بر SVM شده است.

قسمت های تشکیل دهنده این مقاله به قرار زیر است. در بخش ۲، به ذکر مقدمات اولیه مورد نیاز می پردازیم. سپس، در بخش ۳، به معرفی ماشین بردار پشتیبان فازی-استوار و توابع عضویت مختلف می پردازیم. در بخش ۴، روش های پیشنهاد شده را توسط داده های شبیه سازی شده و واقعی مورد آزمایش قرار می دهیم. در انتها، در بخش ۵ به نتیجه گیری و جمع بندی می پردازیم.

۲. مروری بر برخی مقدمات اولیه

^۲Fuzzification

^۳Membership function

^۴Euclidean

^۵Minimum Covariance Determinant (MCD)

^۶Mahalanobis

۱.۲. ماشین بردار پشتیبان. ماشین بردار پشتیبان یک الگوریتم طبقه‌بندی‌کننده دودویی است که هدف آن بیشینه‌سازی حاشیه اطراف مرز جداکننده و کمینه‌سازی خطای آموزشی است. بیشینه کردن حاشیه منجر به افزایش خاصیت تعمیم‌پذیری SVM می‌شود. طراحی یک مسأله طبقه‌بندی‌کننده دودویی با استفاده از ماشین بردار پشتیبان به صورت زیر است. فرض کنید $(\mathbf{x}_1, y_1), (\mathbf{x}_2, y_2), \dots, (\mathbf{x}_n, y_n)$ که $\mathbf{x}_i \in \mathbb{R}^p$ و $y_i \in \{-1, +1\}$ برچسب‌هایی مربوط به هر طبقه باشند. ابر صفحه جداکننده به صورت $\mathbf{w}^T \varphi(\mathbf{x}_i) - b = 0$ تعریف می‌شود که در آن $\mathbf{w}$ بردار عمود بر ابر صفحه و b مشخص‌کننده فاصله ابر صفحه تا مبدأ مختصات است. تابع غیر خطی $\varphi(\mathbf{x}_i)$ برای نگاشت مشاهدات به فضایی با ابعاد بالاتر مورد استفاده قرار می‌گیرد. ابر صفحه بهینه توسط مسأله بهینه‌سازی اولیه به فرم زیر به دست می‌آید.

$$\min_{\mathbf{w}, \xi} \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^n \xi_i \quad (1.2)$$

بطوری‌که،

$$y_i(\mathbf{w}^T \varphi(\mathbf{x}_i) - b) \geq 1 - \xi_i; \quad \xi_i \geq 0; \quad 1 \leq i \leq n.$$

ξ_i متغیرهای “کمکی”^۷ هستند تا اجازه ی جداسازی اشتباه را به نمونه‌هایی که در مکان مناسب قرار نگرفته‌اند،^۸ بدهند. نقش پارامتر C نیز ایجاد تعادل بین داشتن مدلی با حاشیه بزرگ و داشتن کمترین مقدار خطا خواهد بود. شایان ذکر است که، یافتن ابر صفحه بهینه با تبدیل مسأله بهینه‌سازی به فرم ثانویه به صورت ساده‌تری قابل محاسبه خواهد بود. حل این مسأله با استفاده از برقراری شرایط “KKT”^۹ و استفاده از ضرایب لاگرانژ میسر می‌باشد.

$$\max_{\alpha} \sum_{i=1}^n \alpha_i - \frac{1}{2} \sum_{i=1}^n \sum_{j=1}^n \alpha_i \alpha_j y_i y_j K(\mathbf{x}_i, \mathbf{x}_j) \quad (2.2)$$

به طوری‌که،

$$\sum_{i=1}^n \alpha_i y_i = 0; \quad 0 \leq \alpha_i \leq C; \quad 1 \leq i \leq n,$$

که در آن برداری از مقادیر نامنفی ضرایب لاگرانژ است و $K(\mathbf{x}_i, \mathbf{x}_j) = (\varphi(\mathbf{x}_i) \cdot \varphi(\mathbf{x}_j))$ تابع هسته است. با حل مسأله بهینه‌سازی فوق، ضرایب α_i و در نتیجه بردار نرمال ابر صفحه $\mathbf{w} = \sum_{i=1}^n \alpha_i y_i \varphi(\mathbf{x}_i)$ به دست می‌آید. با استناد به شرایط KKT مقدار b نیز قابل محاسبه است. در نتیجه، نمونه‌ی $\mathbf{x}_i$ را می‌توان بر اساس علامت به دست آمده از تابع طبقه‌بندی آن به فرم زیر طبقه‌بندی کرد

$$f(\mathbf{x}) = \text{sign}(\mathbf{w}^T \varphi(\mathbf{x}_i) + b). \quad (3.2)$$

^۷Slack variables

^۸Noisy observations

^۹Karush–Kuhn–Tucker (KKT) conditions

۲.۲. ماشین بردار پشتیبان فازی. تصمیم نهایی در مورد طبقه‌بندی داده‌ها با استفاده از SVM بر مبنای تعداد محدودی از نقاط نمونه‌های آموزشی گرفته می‌شود که به آن‌ها بردارهای پشتیبان گفته می‌شود. اما وجود مشاهده‌ی دورافتاده‌ای در داخل این نمونه‌ها باعث کاهش دقت SVM خواهد شد. اگر چه SVM یک ابزار قدرتمند طبقه‌بندی است اما رویکرد "دقیق"^{۱۰} آن در طبقه‌بندی و یکسان در نظر گرفتن همه نقاط، چه آن‌هایی که از نظر فاصله به مرکز داده‌ها نزدیک‌تر هستند و چه آن‌هایی که در فاصله دورتری نسبت به مرکز داده‌ها قرار گرفته‌اند، باعث شده است که دقت آن تحت تأثیر مشاهدات دورافتاده قرار بگیرد.

در بسیاری مسائل کاربردی به نظر می‌رسد بعضی از نقاط مانند مشاهدات دورافتاده را نمی‌توان به صورت قطعی جزء یک کلاس خاص در نظر گرفت. بنابراین، تفکیک نقاط عادی از مشاهدات دورافتاده در هر کلاس، بسیار حائز اهمیت خواهد بود. فازی‌سازی الگوریتم ماشین بردار پشتیبان برای افزایش قدرت طبقه‌بندی آن مؤثر بوده، که در تحقیقات مختلف مانند [۵]، [۱۱] و [۴] به اثبات رسیده است. هدف SVM فازی نیز مانند SVM، بیشینه‌سازی حاشیه مرز جداکننده است در حالی که، به طور هم‌زمان به کاهش تأثیر مشاهدات دورافتاده نیز می‌پردازد.

۳.۲. فاصله ماهالانوبیس در مقابل فاصله اقلیدسی. این فاصله به عنوان یکی از شناخته‌شده‌ترین فواصل در آمار چند متغیره است که وابستگی بین متغیرها در ساختار آن، در نظر گرفته شده است که به صورت زیر تعریف می‌شود

$$MD = \sqrt{(\mathbf{x} - \bar{\mathbf{x}})^T \mathbf{S}^{-1} (\mathbf{x} - \bar{\mathbf{x}})}. \quad (۴.۲)$$

که در آن $\mathbf{x}$ یک بردار p -بعدی در فضای $\mathbf{R}^p$ ، $\bar{\mathbf{x}}$ میانگین و $\mathbf{S}^{-1}$ معکوس برآوردگر کوواریانس نمونه می‌باشد [۶].

میانگین تجربی به عنوان شناخته‌شده‌ترین معیار برای شناسایی مشاهدات دورافتاده به صورت خیلی زیاد تحت تأثیر قرار گرفته، به نحوی که حتی وجود یک مشاهده دورافتاده باعث ایجاد تغییرات زیادی در این معیار خواهد شد. بنابراین، به منظور شناسایی و رفع این مشکل، استفاده از برآوردگرهای استوار پیشنهاد شده است [۱۰].

۴.۲. برآوردگر مینیمم دترمینان کوواریانس. یکی از پرکاربردترین برآوردگرهای مکان و مقیاس استوار، برآوردگر مینیمم دترمینان کوواریانس است که توسط روسیو [۹] معرفی شده است. هدف این برآوردگر پیدا کردن زیر مجموعه‌ای با بیشترین تعداد مشاهدات است که ماتریس کوواریانس آن دارای کمترین مقدار دترمینان باشد.

$$H_0 = \underset{H}{\operatorname{argmin}} \det(\operatorname{cov}(\mathbf{x}_i | i \in H)).$$

زیر مجموعه‌ی H ، باید حداقل شامل $h = \lfloor \frac{(n+p+1)}{2} \rfloor$ مشاهده باشد ([] نشان دهنده تابع جزء صحیح است). یعنی، در مجموع به تعداد C_h^n زیر مجموعه h عضوی وجود خواهد داشت که باید مورد آزمایش قرار گیرند. برآوردگر مکان و مقیاس MCD، در واقع میانگین و کوواریانس زیر مجموعه H_0 است.

^{۱۰}Crisp

۳. ماشین بردار پشتیبان فازی-استوار

در SVM فازی، عملکرد هر یک از نمونه‌های آموزشی توسط تابع عضویت مشخص می‌شود. در نتیجه، انتخاب یک تابع عضویت مناسب، یکی از مهمترین مراحل در SVM فازی است. معمولاً، تابع عضویت بر مبنای فاصله بین نقاط تا مرکز ساخته می‌شود. تابع عضویت زیر با ویژگی‌هایی که قبلاً ذکر شد، توسط لین و ونگ معرفی شد [۵].

$$\mu_{EUC}(\mathbf{x}_i) = \begin{cases} 1 - \frac{\|\bar{\mathbf{x}}_+ - \mathbf{x}_i\|}{(r_+ + \delta)} & \text{if } \mathbf{x}_i \in C_+ \\ 1 - \frac{\|\bar{\mathbf{x}}_- - \mathbf{x}_i\|}{(r_- + \delta)} & \text{if } \mathbf{x}_i \in C_- \end{cases}$$

که $\bar{\mathbf{x}}_{\pm} = \frac{1}{|I_{\pm}|} \sum_{i \in I_{\pm}} \mathbf{x}_i$ و $r_{\pm} = \max \|\bar{\mathbf{x}}_{\pm} - \mathbf{x}_i\|$ به ترتیب میانگین و "شعاع"^{۱۱} هر کدام از کلاس‌ها هستند.

به عنوان اولین روش پیشنهادی می‌توان با جایگزینی میانگین با یک میانگین استوار مانند MCD، به استوارسازی آن پرداخت. پس در این روش نیز برای محاسبه فاصله، از فاصله اقلیدسی استفاده کردیم. این روش با $\mu_{EUCMCD}(\mathbf{x}_i)$ شناخته می‌شود. استوارسازی شعاع هر کلاس نیز مانند روش فوق با جایگذاری برآوردگر استوار به جای برآوردگر کلاسیک انجام می‌شود. تابع عضویت اقلیدسی استوار به صورت زیر است که آن را با $\mu_{EUCMCD}(\mathbf{x}_i)$ نمایش می‌دهیم.

$$\mu_{EUCMCD}(\mathbf{x}_i) = \begin{cases} 1 - \frac{\|\bar{\mathbf{x}}_{MCD+} - \mathbf{x}_i\|}{(r_+ + \delta)} & \text{if } \mathbf{x}_i \in C_+ \\ 1 - \frac{\|\bar{\mathbf{x}}_{MCD-} - \mathbf{x}_i\|}{(r_- + \delta)} & \text{if } \mathbf{x}_i \in C_- \end{cases}$$

۱.۳. تابع عضویت بر مبنای فاصله ماحالانوبیس. محاسبه فاصله یا میزان مشابهت، اولین مرحله در اکثر روش‌های چند متغیره است. فاصله ماحالانوبیس را می‌توان به عنوان جایگزینی برای فاصله اقلیدسی به حساب آورد که فاصله بین مشاهدات تا مرکز را به صورت معکوس با کوواریانس داده‌ها وزن می‌دهد.

$$\mu_{MAH}(\mathbf{x}_i) = \begin{cases} 1 - \frac{\sqrt{(\mathbf{x}_i - \bar{\mathbf{x}}_+)^T \mathbf{S}_+^{-1} (\mathbf{x}_i - \bar{\mathbf{x}}_+)}}{(r_+ + \delta)} & \text{if } \mathbf{x}_i \in C_+ \\ 1 - \frac{\sqrt{(\mathbf{x}_i - \bar{\mathbf{x}}_-)^T \mathbf{S}_-^{-1} (\mathbf{x}_i - \bar{\mathbf{x}}_-)}}{(r_- + \delta)} & \text{if } \mathbf{x}_i \in C_-, \end{cases}$$

که $r_{\pm} = \max(MD_{i\pm})$ نقطه‌ای با بزرگ‌ترین مقدار فاصله ماحالانوبیس، در هر دو کلاس است. همانطور که عنوان شد وظیفه تابع عضویت، اختصاص درجه اهمیت متناظر به هر یک از نقاط می‌باشد. در این روش به طور مشابه با روش‌های قبل، دورترین نقطه با استفاده از فاصله ماحالانوبیس به دست می‌آید که از آن در تعریف مشاهدات دورافتاده استفاده می‌شود.

۲.۳. تابع عضویت بر مبنای فاصله ماحالانوبیس استوار. به دلیل به کار رفتن میانگین و کوواریانس در ساختار فاصله ماحالانوبیس این فاصله، تحت تأثیر مشاهدات دورافتاده قرار

^{۱۱}Radius

می‌گیرد. بنابراین به‌عنوان یک جایگزین برای آن پیشنهاد می‌شود که از برآوردگر MCD استفاده شود. تابع عضویت بر مبنای فاصله ماهالانوبیس استوار به‌صورت زیر تعریف می‌شود.

$$\mu_{MAHMCD}(\mathbf{x}_i) = \begin{cases} 1 - \frac{\sqrt{(\mathbf{x}_i - \bar{\mathbf{x}}_{MCD+})^T \mathbf{S}_{MCD+}^{-1} (\mathbf{x}_i - \bar{\mathbf{x}}_{MCD+})}}{(r_{MCD+} + \delta)} & \text{if } \mathbf{x}_i \in C_+ \\ 1 - \frac{\sqrt{(\mathbf{x}_i - \bar{\mathbf{x}}_{MCD-})^T \mathbf{S}_{MCD-}^{-1} (\mathbf{x}_i - \bar{\mathbf{x}}_{MCD-})}}{(r_{MCD-} + \delta)} & \text{if } \mathbf{x}_i \in C_- \end{cases}$$

که $\bar{\mathbf{x}}_{MCD\pm}$ و $\mathbf{S}_{MCD\pm}$ برآوردگرهای مکان و مقیاس MCD برای کلاس‌های مثبت و منفی هستند. $r_{\pm} = \max(RMD_{i_{\pm}})$ دورترین مشاهده نسبت به مرکز در هر دو کلاس است. ابر صفحه جداکننده فازی-استوار، بر اساس فرمول کلاسیک SVM به‌دست می‌آید. که در آن فازی‌سازی و استوارسازی در تابع عضویت μ نهفته‌اند. تابع عضویت جدید بایستی در فرمولاسیون SVM گنجانده شود. روش حل مسأله دقیقاً مشابه به حل مسأله SVM کلاسیک است. فرم اولیه مسأله بهینه‌سازی به‌صورت زیر است،

$$\min_{\mathbf{w}, \xi} \frac{1}{2} \|\mathbf{w}\|^2 + C \sum_{i=1}^n \xi_i \mu_i,$$

به‌طوری‌که،

$$y_i(\mathbf{w}^T \boldsymbol{\varphi}(\mathbf{x}_i) - b) \geq 1 - \xi_i; \quad \xi_i \geq 0; \quad 1 \leq i \leq n.$$

توجه داشته باشیم که مضرب $\xi_i \mu_i$ را به‌عنوان معیار خطا نیز می‌توان در نظر گرفت که برای هر نقطه وزن متفاوتی را در نظر می‌گیرد. معیار جدید در واقع یک عامل ترکیبی از دو عبارت ξ_i یعنی، متغیر جریمه و μ_i که نشان‌دهنده رفتار نمونه در هر کلاس است، می‌باشد.

۴. آزمایش با داده‌های شبیه‌سازی شده و واقعی

در این بخش داده‌های شبیه‌سازی شده و واقعی مورد آزمایش قرار می‌گیرند تا کارایی روش‌های پیشنهاد شده، بررسی گردد.

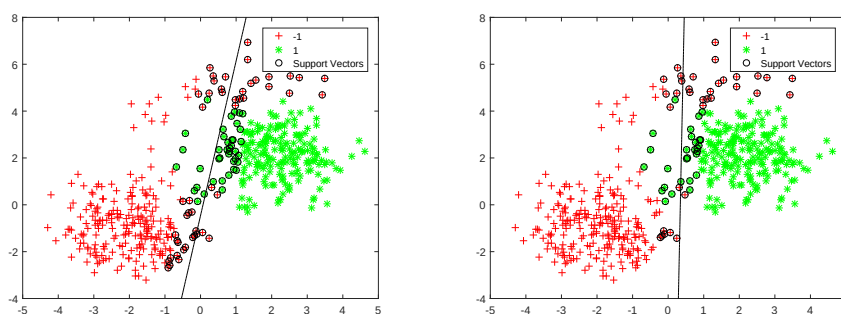
۱.۴. داده‌های شبیه‌سازی شده. در مجموع ۴۰۰ مشاهده از توزیع نرمال دو متغیره $MVN(\boldsymbol{\mu}, \boldsymbol{\Sigma})$ با مشخصات زیر برای هر کلاس تولید شد. کلاس مثبت و منفی دارای میانگین $\boldsymbol{\mu}_+$ و $\boldsymbol{\mu}_-$ هستند. کوواریانس هر دو کلاس برابر با $\boldsymbol{\Sigma}_{\pm}$ است.

$$\boldsymbol{\mu}_+ = \begin{bmatrix} 2 \\ 2 \end{bmatrix} \quad \boldsymbol{\mu}_- = \begin{bmatrix} -2 \\ -1 \end{bmatrix} \quad \text{و} \quad \boldsymbol{\Sigma}_{\pm} = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}.$$

قسمت کوچکی از داده‌ها به‌عنوان مشاهده دورافتاده، از توزیعی متفاوت $MVN(\boldsymbol{\mu}^*, \boldsymbol{\Sigma}^*)$ و برچسب منفی تولید شده‌اند.

$$\boldsymbol{\mu}^* = \begin{bmatrix} 1 \\ 5 \end{bmatrix} \quad \text{و} \quad \boldsymbol{\Sigma}^* = \begin{bmatrix} 2.5 & 1 \\ 1 & 1 \end{bmatrix}.$$

^{۱۲}Robust Mahalanobis distance



شکل ۱: مرز جداکننده تولید شده توسط SVM (پنل سمت چپ) و SVM فازی-استوار (پنل سمت راست)

بعلاوه یک مرکز استوار به کار گرفته شده است، محل قرارگیری مرز جداکننده کمتر تحت تأثیر قرار گرفته است. این امر نیز بر محل قرارگیری مرز جداکننده و در نتیجه دقت SVM تأثیر زیادی دارد.

۲.۴. داده‌های واقعی. در این قسمت، برخی داده‌های واقعی که بسیار در زمینه طبقه‌بندی مورد استفاده قرار می‌گیرند، به کار برده شده است. از جمله آن‌ها می‌توان به Fisher Iris، Pima Indian Diabetes و Wine اشاره کرد.

جدول ۲: دقت روش‌های SVM، FSVM-EUC، RFSVM-EUC، FSVM-MAH و RFSVM-MAH برای داده‌های واقعی و $C = 2^3$.

داده‌ها	SVM	FSVM-EUC	RFSVM-EUC	FSVM-MAH	RFSVM-MAH
Iris	0.9333	0.9333	0.9667	0.9000	0.9000
Wine	0.9743	0.9743	۱	0.9743	0.9743
Diabetes	0.7532	0.7532	0.7792	0.7705	0.7792

کارایی SVM و روش‌های پیشنهادی در جدول ۲ گزارش شده است. جدول فوق نیز به ارجحیت SVM فازی بر اساس فاصله اقلیدسی استوار تأکید دارد. پس از آن، روش‌های SVM فازی بر اساس فاصله مالهالانوبیس و مالهالانوبیس استوار قرار دارند. بالاتر بودن دقت SVM فازی بر اساس فاصله اقلیدسی، را می‌توان دلیل وجود مشاهدات دورافتاده در داده‌ها قلمداد کرد. اما پیشنهاد می‌شود که محقق ابتدا دقت روش‌های مختلف را محاسبه و سپس روشی را برای طبقه‌بندی انتخاب کند که دارای دقت بیشتری باشد.

۵. نتیجه‌گیری

حضور مشاهدات دورافتاده در داده‌ها باعث تأثیر قرار گرفتن کارایی ماشین بردار پشتیبان و در نتیجه کاهش توانایی تعمیم آن می‌شود. بنابراین، در این مقاله با استفاده از انواع مختلف توابع عضویت فازی به مقابله با این مشکل پرداختیم. اولین روش پیشنهادی، استوارسازی تابع عضویت

بر مبنای فاصله اقلیدسی است. علاوه بر آن، فاصله ماهالانوبیس به عنوان جایگزینی برای فاصله اقلیدسی نیز مورد استفاده قرار گرفت. همچنین، برای استوارسازی فاصله ماهالانوبیس در مقابل مشاهدات دورافتاده، برآوردگر مینیم دترمینان کوواریانس به کار گرفته شد. استوارسازی تابع عضویت فازی، به صورت مضاعف باعث کاهش تأثیر مشاهدات دورافتاده بر دقت SVM شد. کارایی SVM و روش های پیشنهادی با استفاده از داده های شبیه سازی شده و واقعی مورد آزمایش قرار گرفت. SVM فازی با استفاده از فاصله اقلیدسی استوار نسبت به سایر روش ها عملکرد بهتری داشت و تابع عضویت بر اساس فاصله ماهالانوبیس استوار در جایگاه دوم قرار گرفت.

مراجع

1. Cortes, C., & Vapnik, V. (1995). *Support-vector networks*. Machine learning, 20(3), 273-297.
2. Hawkins, D. M. (1980). *Identification of outliers* (Vol. 11). London: Chapman and Hall.
3. Lee, Y. (2010). *Support vector machines for classification: a statistical portrait*. Statistical Methods in Molecular Biology, 347-368.
4. Inoue, T., & Abe, S. (2001). *Fuzzy support vector machines for pattern classification*. In Neural Networks, 2001. Proceedings. IJCNN'01. International Joint Conference on (Vol. 2, pp. 1449-1454). IEEE.
5. Lin, C. F., & Wang, S. D. (2002). *Fuzzy support vector machines*. IEEE transactions on neural networks, 13(2), 464-471.
6. Mahalanobis, P. C. (1936). *On the generalized distance in statistics*. Proceedings of the National Institute of Sciences (Calcutta), 2, 49-55.
7. Maronna, R. A. R. D., Martin, R. D., & Yohai, V. (2006). *Robust statistics* (pp. 978-0). John Wiley & Sons, Chichester. ISBN.
8. Maronna, R. A., & Zamar, R. H. (2002). *Robust estimates of location and dispersion for high-dimensional datasets*. Technometrics, 44(4), 307-317.
9. Rousseeuw, P. J. (1984) *Least median of squares regression*. Journal of the American statistical association, 79(388):871-880.
10. Rousseeuw, P. J., & Hubert, M. (2011). *Robust statistics for outlier detection*. Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery, 1(1), 73-79.
11. Tsujinishi, D., & Abe, S. (2003). *Fuzzy least squares support vector machines for multiclass problems*. Neural Networks, 16(5), 785-792.
12. Vapnik, V. N. (1999). *An overview of statistical learning theory*. IEEE transactions on neural networks, 10(5), 988-999.



آزمون فرضیه‌های دقیق بر اساس داده‌های فازی نرمال با رویکرد p -مقدار

محبوبه سادات مدنی^۱ * و محمدرضا ربیعی^۱

^۱ دانشگاه صنعتی شاهرود

mah.madani@mailfa.com

^۲ آدرس نویسنده دوم

rabie1354@yahoo.com

چکیده. در آزمون فرضیه‌های کلاسیک تمام مولفه‌های یک مدل از قبیل داده‌ها، فرضیه‌ها و نحوه بدست آوردن آزمون باید دقیق و بدون ابهام باشند. ولی در عمل و دنیای واقعی برای تعریف مساله به مواردی نادقیق و دارای ابهام برخورد می‌کنیم که این ابهام ممکن است در داده‌ها، فرضیه‌ها و یا هر دو صورت گیرد. بنابراین سه حالت فازی زیر رخ می‌دهد

۱- آزمون فرضیه‌های دقیق با داده‌های نادقیق،

۲- آزمون فرضیه‌های نادقیق با داده‌های دقیق،

۳- آزمون فرضیه‌های نادقیق با داده‌های نادقیق.

در این مقاله مساله آزمون فرضیه‌های دقیق با داده‌های نادقیق نرمال را بر اساس رویکرد p -مقدار مبتنی بر اصل گسترش، مورد بحث قرار می‌دهیم. در این روش از مقایسه p -مقدار فازی و سطح معنی‌داری فازی استفاده شده و برای توضیح بیشتر یک مثال عددی نیز مطرح شده است.

۱. مقدمه

آزمون فرضیه‌های آماری برای تصمیم‌گیری در اکثر مسایل کاربردی و علمی آماری مورد توجه قرار گرفته است. در بسیاری از مسایل آماری ممکن است در فرضیه‌ها یا داده‌ها و یا هر دو ابهام وجود داشته باشد. در این مقاله به بررسی آزمون فرضیه‌های دقیق بر اساس داده‌های فازی نرمال با رویکرد p -مقدار می‌پردازیم. مساله آزمون فرضیه‌ها با داده‌های مبهم توسط کازالس و همکاران [۴، ۵] مورد مطالعه قرار گرفته است. برای اولین بار موضوع آزمون فرضیه‌ها بر پایه داده‌های فازی توسط تاناکا و همکاران [۶] به شکل نظریه تصمیم فازی مطرح شد. فیلموسر و فیتل [۷] مفهوم p -مقدار برای آزمون فرضیه بر پایه داده‌های فازی را معرفی کردند. پرچی و همکاران [۸] مفهوم p -مقدار فازی را برای آزمون فرضیه‌های فازی (داده‌ها دقیق) تعریف و بررسی کرده‌اند. همچنین در مطالعه ای دیگر آزمون فرضیه‌های فازی را بر اساس p -مقدار با مشاهدات فازی مورد بحث قرار داده‌اند [۹]. اخیراً رویکرد کلی مبتنی بر p -مقدار برای آزمون

واژگان کلیدی. آزمون فرضیه، فرضیه‌های فازی، p -مقدار فازی.
* سخنران

کیفیت با در نظر گرفتن فرضیه فازی را مطرح کرده‌اند [۱۰]. برای بررسی برخی دیدگاه‌های دیگر، در آزمون فرضیه‌های آماری در محیط فازی به عارفی و طاهری [۳، ۱۴، ۱۵، ۱۶] مراجعه نمایید. در بخش اول این مقاله تعاریف و مقدمات مورد نیاز را بیان می‌کنیم و در بخش دوم به معرفی p -مقدار فازی، سطح معنی‌داری فازی و بررسی آزمون فرضیه‌ها با داده‌های فازی نرمال می‌پردازیم. سپس مثالی کاربردی برای روشن شدن این روش بکار می‌بریم. و در بخش آخر به نتیجه‌گیری می‌پردازیم.

۲. تعاریف و مفاهیم اولیه

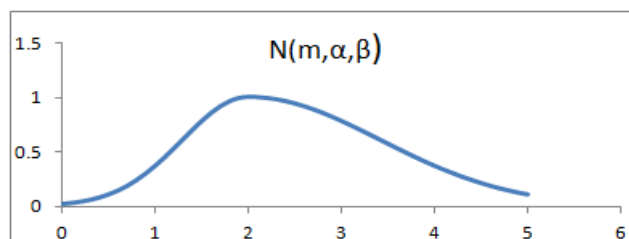
در این بخش برخی تعاریف و مفاهیم اولیه را که در بخش‌های بعد مورد نیاز است، بیان می‌کنیم. δ -برش و تکیه‌گاه یک مجموعه فازی A را به ترتیب با نمادهای A_δ و $\text{supp}(A)$ نشان می‌دهیم. فرض کنید $\mathbb{R}$ مجموعه تمام اعداد حقیقی باشد و

$$F_C(\mathbb{R}) = \{A | A : \mathbb{R} \rightarrow [0, 1], \text{ باشد پیوسته باشد},\}$$

$$F_N(\mathbb{R}) = \{N = (m, \alpha, \beta)_N | m \in \mathbb{R}, \alpha \text{ پهنای چپ باشد}, \beta \text{ پهنای راست و}\}$$

که در آن تابع $N : \mathbb{R} \rightarrow [0, 1]$ نشان دهنده «تقریباً برابر بودن» است و به صورت زیر تعریف می‌شود (شکل ۱)

$$N(m, \alpha, \beta)(x) = \begin{cases} e^{-\left(\frac{m-x}{\alpha}\right)^2} & x \leq m \\ e^{-\left(\frac{x-m}{\beta}\right)^2} & x > m \end{cases}$$



شکل ۱: تابع N برای نمایش برخی مفاهیم فازی نرمال

بنابراین $N(m, 0, 0) = I_{\{m\}}$ که در آن I تابع نشانگر و $m \in \mathbb{R}$ می‌باشد. در ادامه اشاره ای به شیوه آزمون فرضیه‌های آماری، مبتنی بر رویکرد p -مقدار می‌کنیم [۱۱]. در آزمون‌های پارامتری غیرتصادفی، فضای تمامی مقادیر ممکن برای آماره آزمون T به دو ناحیه رد H_0 (ناحیه بحرانی) و ناحیه پذیرش H_0 (مکمل ناحیه بحرانی) تقسیم می‌شود. با توجه به فرضیه‌های H_0 و H_1 غالباً ناحیه بحرانی به یکی از سه شکل زیر می‌باشد:

$$(۱۰۲) \quad \text{الف) } T \leq t_l \quad \text{ب) } T \geq t_r \quad \text{ج) } T \notin (t_1, t_2)$$

که در آن t_l یا t_r ، یا t_1 ، t_2 چندک‌های توزیع T بوده و با مشخص شدن سطح معنی‌داری آزمون مقادیر عددی آن‌ها تعیین می‌گردد.

در آزمون فرضیه‌های معمولی مفهوم p -مقدار با توجه به سه نوع ناحیه بحرانی که در (۱۰۲) مطرح شد به صورت زیر فرمول‌بندی و ساده می‌گردد:

$$\text{الف) } p = P_{\theta_0}(T \leq t) - \text{مقدار}$$

$$\text{ب) } p = P_{\theta_0}(T \geq t) - \text{مقدار}$$

$$\text{ج) } p = 2\min\{P_{\theta_0}(T \leq t), P_{\theta_0}(T \geq t)\} = \begin{cases} 2P_{\theta_0}(T \geq t) & t \geq m \\ 2P_{\theta_0}(T \leq t) & t \leq m \end{cases}$$

که در آن t مقدار مشاهده شده آماره آزمون T ، θ_0 نقطه مرزی فرضیه H_0 و m میانه توزیع T می‌باشد. بنابراین می‌توان گفت که p -مقدار تابعی از:

(۱) نوع فرضیه H_1 ،

(۲) نقطه مرزی فرضیه H_0 ،

(۳) توزیع آماره آزمون و

(۴) مقدار مشاهده شده توزیع آماره آزمون

می‌باشد [۲].

براساس مقایسه p -مقدار مشاهده شده با α یعنی سطح معنی‌داری آزمون، تصمیم در مورد فرضیه‌های مورد آزمون، اتخاذ می‌گردد. اگر p -مقدار کمتر از α باشد، آنگاه فرضیه H_0 در سطح معنی‌داری α رد و در غیر این صورت پذیرفته می‌شود. در این مقاله آزمون فرضیه‌ها در یک محیط فازی در نظر گرفته شده‌اند و ابهام موجود در داده‌ها به p -مقدار منتقل شده و p -مقدار به صورت مجموعه فازی p بدست آمده است.

۳. آزمون فرضیه‌های دقیق براساس داده‌های فازی نرمال

فرض کنید مشاهدات استفاده شده در آزمون فرضیه دقیق H_0 در مقابل فرضیه دقیق H_1 ، دارای ابهام بوده و به صورت $\tilde{x}_1, \dots, \tilde{x}_n \in F_c(R)$ گزارش شده باشند. در این حالت با استفاده از اصل گسترش، می‌توان گفت که مقدار مشاهده شده آماره آزمون، یک مجموعه فازی به صورت $t = h(\tilde{x}_1, \dots, \tilde{x}_n)$ است که تابع عضویت آن نیز قابل محاسبه می‌باشد [۱، ۷]. در این حالت p -مقدار به عنوان تابعی از t خود نیز به یک مجموعه فازی تبدیل می‌شود که آن را به اختصار با P نشان می‌دهیم. در قضیه ۱۰۳، δ -برش‌های p -مقدار فازی برای سه نوع ناحیه بحرانی مطرح شده در (۱۰۲)، ارایه شده است.

قضیه ۱۰۳. فرض کنید f_θ ، $X \sim \theta$ ، $\theta \in \Theta$. در مساله آزمون فرضیه‌ها با داده‌های فازی، δ -برش‌های p -مقدار فازی در سه حالت مطرح شده در ۱۰۲ به صورت زیر می‌باشد [۲، ۷].

$$\text{الف) } P_\delta = [P_{\theta_0}(T \leq t_1(\delta)), P_{\theta_0}(T \leq t_2(\delta))]$$

$$\text{ب) } P_\delta = [P_{\theta_0}(T \leq t_1(\delta)), P_{\theta_0}(T \leq t_2(\delta))]$$

$$P_{\delta} = \begin{cases} [2P_{\theta_0}(T \geq t_2(\delta), 2P_{\theta_0}(T \geq t_1(\delta))] & \text{if } t_l \geq m \\ [2P_{\theta_0}(T \leq t_2(\delta), 2P_{\theta_0}(T \leq t_1(\delta))] & \text{if } t_l \leq m \end{cases} \quad (ج)$$

که در آن $\theta_0, \delta \in (0, 1]$ نقطه مرزی فرضیه H_0 و توابع t_2, t_1 با استفاده از رابطه

$$t_{\delta} = [t_1(\delta), t_2(\delta)]$$

حاصل می‌شوند. در حالت (ج) داریم

$$t_l = \inf\{t | n \in \text{supp}(t)\}, t_r = \sup\{t | t \in \text{supp}(t)\}$$

نکته ۲.۳. باتوجه به قضیه ۱.۳ در آزمون‌های دوطرفه با یک مساله تصمیم‌گیری سه تصمیمه روبرو هستیم. به عبارت دیگر در حالت (ج)، یکی از سه تصمیم زیر را اتخاذ می‌کنیم

- (۱) فرضیه صفر پذیرش و فرضیه مقابل رد شود،
- (۲) فرضیه صفر رد و فرضیه مقابل پذیرش شود،
- (۳) هیچ‌کدام از فرضیه‌های صفر و مقابل پذیرش یا رد نشوند.

باتوجه به قضیه ۱.۳، تصمیم ۳ وقتی اتخاذ می‌گردد که ابهام داده‌ها آن قدر زیاد باشد تا به نتیجه $m \in \text{supp}(t)$ منجر شده و این موضوع باعث می‌شود که نتوانیم تصمیم به رد یا پذیرش فرضیه صفر بگیریم.

از طرفی می‌توانیم از مفهوم سطح معنی‌داری فازی به جای سطح معنی‌داری دقیق استفاده کنیم.

تعریف ۳.۳. سطح معنی‌داری فازی، زیرمجموعه‌ای فازی است که بر بازه (۱ و ۰) تعریف می‌شود [۱۲]. لذا مطابق این تعریف می‌توان سطح معنی‌داری غیرفازی و دقیق را به کمک تابع نشانگر $I_{\{m\}}$ به‌طوری‌که $m \in (0, 1)$ نشان داد.

اگر p -مقدار فازی از S کمتر باشد آنگاه فرضیه صفر را در سطح معنی‌داری S رد و در غیر این صورت پذیرفته می‌شود. اما برخلاف آزمون فرضیه‌های معمولی، امکان وجود اشتراک بین مجموعه‌های $\text{supp}(S)$ و $\text{supp}(P)$ وجود داشته و در چنین مواردی نمی‌توان با اطمینان کامل ادعا کرد که $P \succ S$ یا $S \succ P$. بنابراین برای انجام این مقایسه نیاز به معیاری جهت مقایسه دقیق‌تر دو مجموعه فازی S و P داریم که در تعریف ۴.۳ مطرح شده است. در این مقاله از معیار و رابطه D استفاده خواهد شد. لازم به ذکر است که رابطه D دارای ویژگی‌های جابجایی، انتقال‌پذیری و پایایی بر روی $F(\mathbb{R}) \times F(\mathbb{R})$ می‌باشد که در [۱۳] مطرح و اثبات شده‌اند.

تعریف ۴.۳. فرض کنید $A, B \in F_C(\mathbb{R})$ و

$$\Delta_{AB} = \int_{A_{\delta}^U > B_{\delta}^L} (A_{\delta}^U - B_{\delta}^L) d\delta + \int_{A_{\delta}^L > B_{\delta}^U} (A_{\delta}^L - B_{\delta}^U) d\delta \quad (۱.۳)$$

که در آن $\delta \in (0, 1)$ و $A_{\delta}^U = \sup\{x | x \in A_{\delta}\}$ و $A_{\delta}^L = \inf\{x | x \in A_{\delta}\}$ ”درجه بزرگی A از B ” [۲] به‌صورت زیر تعریف می‌گردد:

$$D(A \succ B) = \frac{\Delta_{AB}}{\Delta_{AB} + \Delta_{BA}} \quad (2.3)$$

تعریف ۵.۳. اگر در یک مساله آزمون فرضیه تصمیم به رد یا قبول فرضیه صفر بگیریم، آنگاه $D(P \succ S)$ درجه پذیرش فرضیه صفر، و $D(S \succ P) = 1 - D(P \succ S)$ درجه رد فرضیه صفر نامیده می‌شود.

مثال ۶.۳. فرض کنید که طول عمر مفید باطری‌های تولیدی در یک کارخانه دارای توزیع نرمال با میانگین $\mu = 1300$ و انحراف معیار $\sigma = 30$ ساعت می‌باشد. بعد از اعمال یک سری تصمیمات در خط تولید این کارخانه، نمونه‌ای تصادفی به حجم $n = 10$ لایمپ از خط تولید گرفته شده و از آنجا که اندازه طول عمر مفید باطری یک اندازه نادقیق می‌باشد، اعداد فازی نرمال زیر بعنوان مشاهدات نمونه ثبت شده‌اند.

جدول ۱: اعداد فازی نرمال مثال ۶.۳

$(1261, 4, 17)_N$	$(1287, 36, 15)_N$
$(1346, 31, 26)_N$	$(1330, 24, 18)_N$
$(1329, 31, 26)_N$	$(1301, 13, 19)_N$
$(1317, 19, 16)_N$	$(1269, 28, 15)_N$
$(1353, 28, 16)_N$	$(1337, 36, 18)_N$

حال فرض کنید پرسشی مطرح شده است که آیا طول عمر مفید باطری‌های تولیدی پس از اعمال تصحیحات بر خط تولید افزایش یافته است؟ به دیگر سخن قصد داریم این فرضیه آماری $H_0: \mu = 1300$ را در مقابل $H_1: \mu > 1300$ آزمون کنیم. آماره این آزمون یعنی $T = \bar{X} = \frac{1}{10} \sum_{i=1}^{10} X_i$ در نقطه مرزی فرضیه H_0 دارای توزیع $N(1300, \frac{30^2}{10})$ بوده و مقدار مشاهده شده آن به کمک اصل گسترش به صورت عدد نرمال $(1313, 25, 18)_N$ به دست می‌آید. از آنجا که ناحیه بحرانی از نوع (ب) می‌باشد، با توجه به قضیه ۱۰.۳ -برش‌های p -مقدار فازی به ازای $\delta \in (0, 1]$ به صورت زیر محاسبه می‌شوند:

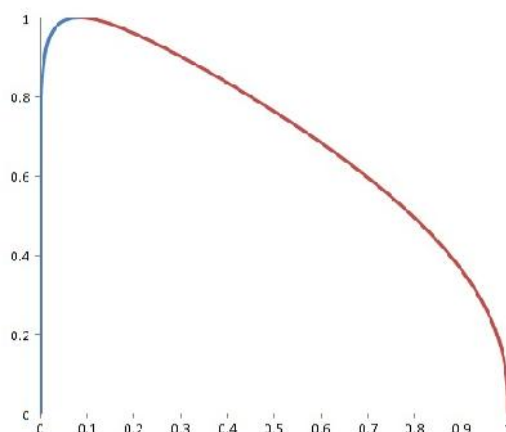
$$\mathbf{P}_\delta = \left[\int_{\frac{t_2(\delta)-1300}{\frac{30}{\sqrt{10}}}}^{\infty} (2\pi)^{-\frac{1}{2}} \exp\left(-\frac{z^2}{2}\right) dz, \int_{\frac{t_1(\delta)-1300}{\frac{30}{\sqrt{10}}}}^{\infty} (2\pi)^{-\frac{1}{2}} \exp\left(-\frac{z^2}{2}\right) dz \right]$$

که در آن $n = (1313, 25, 18)_N$ و در نتیجه نقاط ابتدا و انتهای δ -برش n به صورت

$$t_1(\delta) = 1313 - 25\sqrt{-Ln\delta}, t_2(\delta) = 1313 + 25\sqrt{-Ln\delta}$$

محاسبه می‌شوند. با مشخص شدن کلیه δ -برش‌ها، p -مقدار فازی مشخص می‌شود. نمودار تابع عضویت مربوط در شکل ۲ رسم شده است. در این حالت اگر بخواهیم فرضیه‌های مطرح شده را در سطح معنی‌داری تقریباً ۰.۰۵ (که در اینجا با $S = (0.05, 0.05, 0.05)_N$ مشخص شده) آزمون کنیم، آنگاه $\Delta_{ps} = 0.712$ و $\Delta_{sp} = 0.090$ و در نتیجه فرضیه H_0 با درجه ۰.۸۸۶

پذیرفته خواهد شد.



شکل ۲: تابع عضویت p -مقدار فازی در مثال ۶.۳.

۴. نتیجه‌گیری

در واقع موضوع p -مقدار فازی از آن جهت قابل اهمیت است که راه حل ساده‌تری برای تصمیم در مورد آزمون فرض در اختیار ما قرار می‌دهد، که با مقایسه این شاخص با سطح معنی‌داری فازی در مورد فرض صفر تصمیم اتخاذ می‌نماییم. برای تصمیم‌گیری در مساله آزمون فرضیه‌ها با داده‌های فازی رویکرد p -مقدار که منجر به p -مقدار فازی و سطح معنی‌داری فازی و در نتیجه تصمیم‌گیری فازی شد، رهیافتی کاربردی است.

در این مقاله سعی شده است که رویکرد مطرح شده در [۲]، به وسیله یک مثال عددی با داده‌های فازی نرمال مورد بررسی و بحث قرار گیرد. علاوه بر بسیاری از مسایل کاربردی در علوم مختلف که آمیخته با ابهام می‌باشند، آنالیز واریانس با داده‌های فازی، آزمون فرضیه‌ها برای پارامترهای مدل رگرسیون فازی، و همچنین آزمون فرضیه‌ها بر پایه متغیرهای تصادفی فازی، از موضوعات بالقوه برای تحقیق در آینده بر اساس روش p -مقدار فازی می‌باشند.

مراجع

۱. م. فضلعلی‌پور میان‌دوآب. (۱۳۸۶)، p -مقدار فازی در آزمون فرضیه‌های آماری، پایان‌نامه کارشناسی ارشد آمار، دانشگاه شهید باهنر کرمان.
۲. ع. پرچمی، م. ماشین‌چی. (۱۳۹۰)، آزمون فرضیه در محیط فازی بر پایه p -مقدار، سری سیستم‌های فازی و محاسبات نرم، ۱۸-۱.
۳. م. عارفی، س.م. طاهری. (۱۳۸۶)، آزمون فرض فازی بر پایه آماره آزمون فازی، گزارش هفتمین کنفرانس سیستم‌های فازی و هوشمند، دانشگاه فردوسی مشهد، ۵۷۹-۵۷۵.

4. Casals, M.R., Gil, M.A., and Gil, P. (1986a), *On the use of Zadeh's probabilistic definition for testing statistical hypotheses from fuzzy information*, Fuzzy Sets and Systems. **20**, 175-190.
5. Casals, R., Gil, M.A., and Gil, P. (1986b), *The fuzzy decision problem: an approach to the problem of testing statistical hypotheses with fuzzy information*, European J. of Operational Researches. **27**, 371-382.
6. Tanaka, H., Okuda, T., Asai, K. (1979), *Fuzzy information and decision in statistical model*, In: Gupta, M.M., et al. (Eds.), *Advances in Fuzzy Sets Theory and Applications*, North-Holland, Amsterdam, pp. 303-320.
7. Filzmoser, P., Viertl, R. (2004), *Testing hypothesis with fuzzy data: the fuzzy P-value*, Metrika. **59**, 21-29.
8. Parchami, A., Taheri, S.M., Mashinchi, M. (2008), *Fuzzy p-value in testing hypotheses with crisp data*, Statistical Papers (to appear). **51**, 209-226.
9. Parchami, A., Taheri, S.M., Mashinchi, M. (2012), *Testing fuzzy hypotheses based on vague observations: a p-value approach*, Stat. Papers. **53**, 469-484.
10. Parchami, A., Sadeghpour Gildeh, B., Taheri, S.M., Mashinchi, M. (2017), *A general p-value-based approach for testing quality by considering fuzzy hypotheses*, Journal of Intelligent & Fuzzy Systems. **32**, 1649-1658.
11. Casella, G., Berger, R.L. (1990), *Statistical inference*, Brooks/Cole Publishing, CA.
12. Hoien, M. (2004), *Fuzzy hypotheses testing in a framework of fuzzy logic*, Fuzzy Sets and Systems. **145**, 229-252.
13. Yuan, Y. (1991), *Criteria for evaluating fuzzy ranking methods*, Fuzzy Sets and Systems **43**, 139-157.
14. Arefi, M., Taheri, S. M. (2007), *Testing statistical hypothesis using fuzzy critical region in model with nuisance parameter*, ISI Bulletin, CPM 077, Lisboa, Portugal, pp. 4537-4540.
15. Arefi, M., Taheri, S. M. (2011), *Testing fuzzy hypotheses using fuzzy data based on fuzzy test statistic*, Journal of Uncertain Systems. **5**, pp. 45-61.
16. Arefi, M., Taheri, S. M. (2013), *A new approach for testing fuzzy hypotheses based on fuzzy data*, International Journal of Computational Intelligence Systems. **6**, pp. 318-327.



نمودار تعداد اقلام معیوب با استفاده از درجه عدم انطباق

میترا مهدی پور^۱ * و محمدرضا ربیعی^۲

^۱ دانشگاه صنعتی شاهرود، دانشکده علوم ریاضی، گروه آمار
m21142@gmail.com

^۲ دانشگاه صنعتی شاهرود، دانشکده علوم ریاضی، گروه آمار
rabie.mr87@gmail.com

چکیده. کنترل فرآیند آماری مجموعه‌ای قدرتمند و توانا از ابزارهایی است که با کاهش تغییرپذیری، ثبات و بهبود در کارایی فرآیند را به همراه دارد. نمودارهای کنترل یکی از ابزارهای هفت‌گانه کنترل فرآیند آماری محسوب می‌شوند. نمودارهای کنترل کلاسیک دو حالت محدود صفر و یک را در بر می‌گیرند که اطلاعات کامل و مناسبی در مورد فرآیند تولید در اختیار کاربر قرار نمی‌دهند به همین دلیل در این مقاله ابتدا با استفاده از کیفیت فازی به تعمیم نمودار تعداد اقلام معیوب در حالت فازی پرداخته می‌شود. استفاده از نمودارهای کنترل فازی در مواقعی که درجه بی‌کیفیتی اقلام معیوب برای کاربر مساوی نباشد، مناسب است. در ادامه با یک مثال کاربردی به توضیح این روش پرداخته و با معیارهایی مانند ARL پاسخ ارزشمندتر این نمودارها در مورد میانگین و واریانس در زمان ایجاد تغییر در فرآیند، در مقابل حالت کلاسیک مقایسه می‌شود.

۱. مقدمه

از روش‌های کنترل فرآیند در حین تولید، نمودارهای کنترل است که می‌توان به عنوان ابزاری مؤثر جهت کاهش تغییرپذیری فرآیند استفاده نمود. معادله عمومی نمودارهای کنترل شوهارت به صورت:

$$LCL = \mu_W - k\sigma_W \quad CL = \mu_W \quad UCL = \mu_W + k\sigma_W$$

که در آن، W آماره‌ای که مشخصه کیفی مورد نظر را اندازه‌گیری می‌کند و k فاصله حدود کنترل از خط مرکز برحسب واحد انحراف معیار است [۲]. نمودارهای کنترل را می‌توان به دو گروه کلی تقسیم کرد. اگر بتوان مشخصه کیفیت را اندازه‌گیری و آن را به صورت عددی در مقیاس پیوسته ارائه کرد، آن را یک متغیر کمی می‌نامند. اگر مشخصه کیفیت مورد نظر را نتوان در مقیاس کمی اندازه‌گیری کرد در چنین شرایطی ممکن است بتوان هر محصول را براساس معیار انطباق یا عدم انطباق و یا شمارش تعداد نقص‌های آن گروه بندی نمود. می‌توان دریافت که مقادیر یک متغیر کمی در مقابل مقادیر مشاهده شده از متغیر کیفی که به کمک یک مقیاس کیفی اندازه‌گیری

واژگان کلیدی. کیفیت فازی، نمودار کنترل کیفیت آماری، تعداد اقلام معیوب، منحنی مشخصه عملکرد.
* سخنران

شده، حاوی اطلاعات دقیق‌تر و کامل‌تر از متغیر مربوطه می‌باشد. فقط زیرا در بسیاری از مواقع متغیرهای با مقیاس کیفی محدود به یک تابع عضویت دودویی می‌باشند و در کیفیت فازی برای رفع این محدودیت و برای دسترسی به اطلاعات دقیق‌تر مقادیر اصلی مشخصه کیفی را در نظر گرفته و سپس رفتار مشخصه کیفیت را در مقیاس کمی با استفاده تابع عضویت مناسب به عنوان کیفیت فازی بیان کرد.

بعد از معرفی نظریه مجموعه‌های فازی توسط پروفیسور لطفی زاده در سال ۱۹۶۵ [۱۰]، تعداد زیادی از نویسندگان در حوزه ساختار نمودارهای کنترل، با استفاده از داده‌های فازی فعالیت نموده‌اند. نمودارهای کنترل فازی ابتدا در سال ۱۹۹۰ میلادی توسط وانگ و راز ارائه گردیدند. وانگ و راز [۹] و راز و وانگ [۸] دو رویکرد احتمالی و امکانی را پیشنهاد دادند. کهرمان [۷] با استفاده از اعداد فازی مثلثی الگوهای غیر طبیعی نمودارهای کنترل را آزمون کرد. برش‌های آلفا برای نمودارهای کنترل فازی توسط گولبای و کهرمان [۶] توسعه داده شد و همچنین روش مستقیم فازی [۵] را بدون روش فازی‌زدایی برای ساخت نمودارهای کیفیت استفاده کردند. امیرزاده و همکاران [۳] نمودار p را با استفاده از درجه انطباق و عدم انطباق معرفی کردند. فراز و مقدم [۴] برای متغیرهایی تصادفی و مبهم نمودارهای کنترل را با استفاده از ناحیه پذیرش فازی مطرح کردند.

در این مقاله، بخش ۲ به کیفیت فازی پرداخته می‌شود. بخش ۳ نمودارهای کنترل کیفیت فازی برای تعداد و متوسط درجه عدم انطباق معرفی می‌شود و بخش آخر شامل بیان ساختار این روش با یک مثال و تحلیل و مقایسه حالت فازی و کلاسیک می‌باشد.

۲. کیفیت فازی

حدود مشخصه فنی همان پراکندگی مجاز است که توسط مهندسان ساخت، مشتری و یا طراحان محصول برای رسیدن به هدف مشخص تعیین می‌شود. حدود مشخصه فنی بالا و پایین که به ترتیب با USL و LSL نشان داده می‌شوند. اگر مقدار مشخصه کیفیت اندازه گیری شده یک محصول تولیدی با استانداردها و حدود مشخصه فنی تطابق داشته باشد (متعلق به حدود مشخصه فنی باشد) آن‌گاه آن محصول سالم و در غیر این صورت محصول معیوب شناخته می‌شود. به این ترتیب می‌توان درجه کیفیت و درجه بی کیفیتی مقدار مشخصه کیفیت x را به ترتیب با توابع $C(x)$ و $1 - C(x)$ نشان داد که در آنها

$$C(x) = \begin{cases} 1 & x \in [LSL, USL] \\ 0 & x \notin [LSL, USL] \end{cases}$$

برای یک محصول با حدود مشخصه معلوم، $x \in [LSL, USL]$ ، به معنای با کیفیت بودن آن کالا است و در مقابل $I_{[LSL, USL]}(x) = 0$ به معنای بی کیفیتی آن کالا می‌باشد.

تعریف ۱.۲. [۱] تابع عضویت کیفیت فازی $\tilde{C}$ تابعی از χ به بازه $[0,1]$ است به طوری که اگر x مشخصه کیفیت مشاهده شده برای یک کالا باشد آن‌گاه $\tilde{C}(x)$ میزان کیفیت فازی آن کالا را نشان دهد. به عبارت دیگر، $\tilde{C} \in F(\chi)$ کیفیت فازی برای مشخصه کیفیت X است اگر و تنها اگر $\tilde{C}$ نشان دهنده میزان کیفیت آن کالا به ازای هر $x \in \chi$ باشد که در آن مجموعه $F(\chi)$ تمام مجموعه های فازی روی χ است.

در این مقاله کیفیت فازی به صورت دوزنقه‌ای و با تابع عضویت زیر در نظر گرفته شده است:

$$T(a, b, c, d)(x) = \begin{cases} 0 & x < a \\ \frac{x-a}{b-a} & a \leq x < b \\ 1 & b \leq x \leq c \\ \frac{d-x}{d-c} & c < x < d \\ 0 & x \geq d \end{cases}$$

است که در آن $-\infty < a \leq b \leq c \leq d < \infty$ و در حالت خاص اگر $b = c$ منجر به عدد فازی مثلثی می‌شود. در کیفیت فازی با استفاده از مجموعه‌های فازی، میزان مرغوبیت محصول را از نظر کیفیت درجه بندی می‌کنند که به این ترتیب حالت‌های بیشتری نسبت به دو حالت محدود صفر و یک (انطباق و عدم انطباق) می‌توان برای کیفیت یک محصول در نظر گرفت. با استفاده از این روش می‌توان در مورد روند فرآیند مناسب‌تر قضاوت نمود.

۳. نمودارهای کنترل کیفیت فازی

با فرض اینکه X مشخصه کیفیت محصول دارای توزیع $N(\mu, \sigma^2)$ است حال $X_1, X_2, \dots, X_n$ یک نمونه تصادفی n تایی از مشخصه کیفی X باشد، آن گاه آنچه که در مورد کیفیت غیر فازی در بخش قبل معرفی شد، $N(X_1), N(X_2), \dots, N(X_n)$ یک نمونه تصادفی n تایی از متغیر تصادفی $N(X)$ است به طوریکه به ازای $i = 1, \dots, n$ داریم $N(X_i) \in \{0, 1\}$ در این صورت می‌توان تعداد بی‌کیفیتی‌ها (مثلاً تعداد نقص‌ها، یا تعداد اقلام معیوب) در یک نمونه تصادفی n تایی را به صورت مجموع n عدد صفر یا یک یعنی $\sum_{i=1}^n N(X_i)$ نشان داد. آماره تعمیم یافته‌ی تعداد بی‌کیفیتی‌ها (W) در حالتی که کیفیت به وسیله مجموعه فازی $\tilde{C} = 1 - \tilde{N}$ که متغیر تصادفی $\tilde{N}(X_i) \in [0, 1]$ میزان و درجه‌ی بی‌کیفیتی i امین عضو از نمونه تصادفی را نشان می‌دهند به صورت زیر است:

$$\tilde{W} = \sum_{j=1}^n \tilde{N}(X_j) = \sum_{j=1}^n (1 - \tilde{C}(X_j))$$

که این آماره مجموع درجات بی‌کیفیتی‌ها در نمونه n تایی است. توجه داشته باشید که تعداد بی‌کیفیتی‌ها با نماد w_i و مجموع درجات بی‌کیفیتی‌ها با نماد $\tilde{w}_i \in \mathbb{R}^+ \cup \{0\}$ نشان داده شده‌اند.

حال با توجه به تابع عضویت دوزنقه‌ای داریم:

۱.۳. نمودار کنترل برای متوسط درجه عدم انطباق. [۱] تعداد اقلام معیوب را از یک زمان به زمان دیگر، یا از یک نمونه به نمونه دیگر، می‌توان بر روی نمودار کنترل $\tilde{p}$ و بر اساس کیفیت فازی رسم کرد. تنها وقتی استفاده از نمودار کنترل $\tilde{p}$ به جای نمودار کنترل p در عمل توجیه دارد که درجه‌ی بی‌کیفیتی تمامی اقلام معیوب از دیدگاه کاربر مساوی نباشد. برای رسم نمودار بر اساس اطلاعات m نمونه n تایی، کافیس نقاط $(i, \frac{\tilde{w}_i}{n})$ را به ازای $i = 1, \dots, m$ بر روی

یک نمودار دو بعدی رسم کرده به طوریکه

$$\tilde{w}_i = \sum_{j=1}^n \tilde{N}(x_{ij}) = \sum_{j=1}^n (1 - \tilde{C}(x_{ij}))$$

مجموع درجات بی کیفیتی در نمونه i ام، مقدار مشخصه کیفیت مربوط به j امین عضو نمونه i ام و $\tilde{C}$ تابع عضویت کیفیت فازی است [۱]. اگر تمامی نقاط رسم شده بین حدود کنترل قرار داشته و هیچ گونه رفتار سیستماتیکی از خود نشان ندهند، آنگاه می توان تحت کنترل بودن آماره مربوطه را نتیجه گرفت. اما برای ترسیم نمودار کنترل شوهارت و تشخیص حدود کنترل تنها کافیسیت مقادیر میانگین و انحراف معیار $\tilde{C}(X)$ براساس داده ها برآورد کرد که این کار توسط روابط

$$\bar{\tilde{C}} = \sum_{i=1}^m \frac{\tilde{C}}{m}, \quad s_{\tilde{C}} = \sqrt{\sum_{i=1}^m \frac{s_i^2}{m}}$$

لذا برآورد حدود کنترل نمودار $\tilde{p}$ برابرند با

$$\widehat{LCL} = (1 - \bar{\tilde{C}} - k \frac{s_{\tilde{C}}}{\sqrt{n}})$$

$$\widehat{CL} = 1 - \bar{\tilde{C}}$$

$$\widehat{UCL} = (1 - \bar{\tilde{C}} + k \frac{s_{\tilde{C}}}{\sqrt{n}})$$

۲.۳. نمودار کنترل فازی تعداد اقلام معیوب. برای رسم نمودار براساس اطلاعات m نمونه n تایی، کافیسیت نقاط $(i, \tilde{w}_i)$ را به ازای $m, 1, \dots, i$ بر روی یک نمودار دو بعدی رسم کرد. برای محاسبه حدود کنترل نمودار فازی تعداد اقلام معیوب کافیسیت تا حدود کنترل $\tilde{np}$ را ضرب در حجم نمونه یعنی n کنیم. لذا برآورد حدود کنترل نمودار $\tilde{np}$ برابرند با

$$\widehat{LCL} = n(1 - \bar{\tilde{C}} - k \frac{s_{\tilde{C}}}{\sqrt{n}})$$

$$\widehat{CL} = n(1 - \bar{\tilde{C}})$$

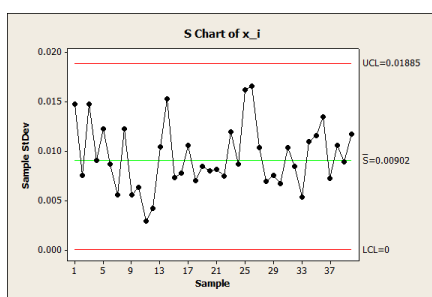
$$\widehat{UCL} = n(1 - \bar{\tilde{C}} + k \frac{s_{\tilde{C}}}{\sqrt{n}})$$

۴. مثال کاربردی

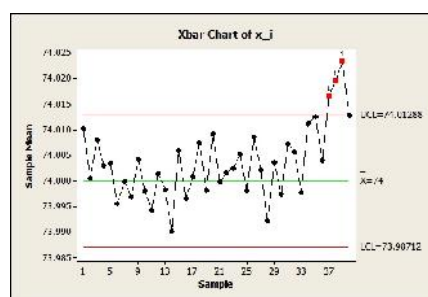
رینگهای پیستون موتور اتومبیلی در یک کارخانه به وسیله یک فرآیند آهنگری تولید می شود. به منظور کنترل قطر داخلی رینگهای تولید شده به وسیله این فرآیند، در نظر داریم از نمودارهای $\bar{x}$ و R استفاده کنیم. بدین منظور تعداد بیست و پنج نمونه که هر کدام شامل پنج رینگ است زمانی که تصور می شود فرآیند تحت کنترل قرار دارد تهیه می گردد. زمانی که میانگین های نمونه ها

و انحراف معیار بر روی این نمودارها رسم می‌شوند هیچگونه اثری از حالت خارج کنترل بودن مشاهده نمی‌گردد. با توجه به حالت کنترل فرآیند با استفاده از نمونه‌های جمع آوری شده، میانگین و انحراف معیار فرآیند تولید را به صورت $\hat{\mu} = \bar{\bar{x}} = 74$ ، $\hat{\sigma} = \frac{\bar{s}}{c_4}$ تخمین می‌زنیم. سپس پانزده نمونه دیگر از فرآیند تولید رینگهای پیستون جمع آوری شده است. حال با توجه به اینکه مشخصات فنی قطر داخلی رینگ پیستون در حالت غیر فازی 74 ± 0.5 است، برای درجه انطباق تابع عضویت مثلثی زیر را در نظر می‌گیریم.

$$\tilde{C}(x) = \begin{cases} 0 & x < 73.95 \\ \frac{x-73.95}{0.05} & 73.95 \leq x \leq 74 \\ \frac{74.05-x}{0.05} & 74 < x \leq 74.05 \\ 0 & x > 74.05 \end{cases}$$

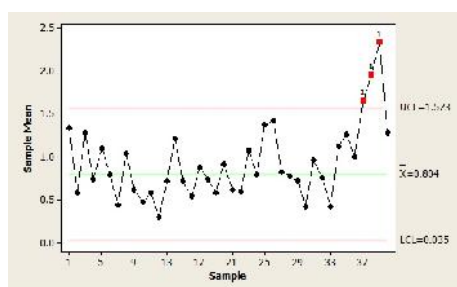


(ب) نمودار s



(آ) نمودار $\bar{x}$

شکل ۱: نمودار $\bar{x}$ و نمودار s برای ۴۰ نمونه ۵ تایی



شکل ۲: نمودار $\bar{np}$ برای ۴۰ نمونه ۵ تایی

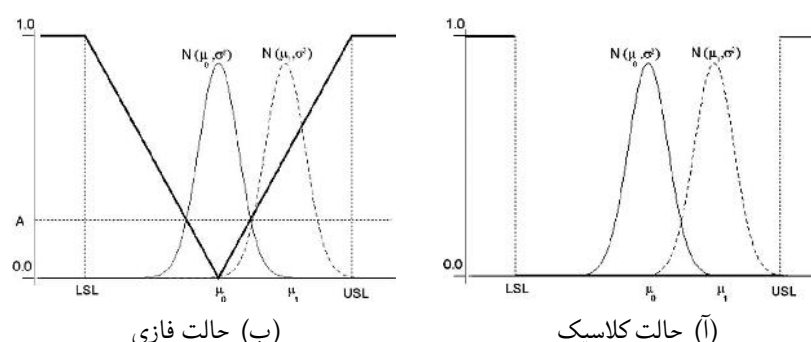
شکل ۳ (آ)، نمودار $\bar{x}$ مانند شکل ۲ خارج از کنترل شدن فرآیند بر اثر تغییر میانگین را نشان می‌دهند. در شکل ۴ نمودار $\bar{np}$ برای $\sum_{i=1}^5 \tilde{N}(x_i)$ با حدود کنترل $CL = 0.804$, $UCL = 1.573$, $LSL = 0$

رسم شده است. درواقع نمودار $\tilde{np}$ بر خلاف نمودار np خارج از کنترل رفتن فرآیند را بر اثر تغییر تشخیص خواهد داد.
در جدول ۱ بعد، درجه عدم انطباق هر محصول با استفاده از تابع عضویت به دست آمده است. [۱]

جدول ۱: درجه عدم انطباق مشاهدات

sample	$\tilde{N}(x_1)$	$\tilde{N}(x_2)$	$\tilde{N}(x_3)$	$\tilde{N}(x_4)$	$\tilde{N}(x_5)$	$\sum_{i=1}^5 \tilde{N}(x_i)$
1	0.6	0.04	0.38	0.16	0.16	1.34
2	0.1	0.16	0.02	0.22	0.08	0.58
3	0.24	0.48	0.42	0.1	0.04	1.28
4	0.04	0.08	0.14	0.3	0.18	0.74
5	0.16	0.14	0.3	0.22	0.28	1.1

۱.۴. مقایسه نمودار کنترل np و نمودار کنترل $\tilde{np}$. همانطور در شکل مشخص است، درجه عدم انطباق محصولات ما در حالت کنترل ($\mu = \mu_0$) و حتی زمانی که میانگین فرآیند به μ_1 تغییر می‌کند، برابر صفر است، به عبارت دیگر نمودار np به این تغییرات عکس العمل نشان نخواهد داد. در صورتی که باتوجه به شکل تعداد محصولات دارای درجه عدم انطباق فازی کمتر از A در حالت کنترل فرآیند و زمانی که میانگین فرآیند تغییر می‌کند، متفاوت است و این بیانگر حساسیت بالاتر نمودار $\tilde{np}$ نسبت به حالت کلاسیک می‌باشد.
منحنی OC معیاری برای ارزیابی حساسیت نمودار کنترل و یا به عبارت دیگر پی بردن به تغییر در تعداد اقلام معیوب فرآیند است و نموداری است که احتمال پذیرش اشتباهی، فرضیه تحت کنترل آماری بودن فرآیند را برحسب تعداد اقلام معیوب نشان می‌دهد.



شکل ۳: مقایسه درجه عدم انطباق برای حالت کلاسیک و فازی با فرآیند تولید نرمال

حال با استفاده از معیارهای OC با استفاده از فرضیه‌های زیر به مقایسه دو نمودار $np, \widetilde{np}$ $chart$ پرداخته خواهد شد.

$$\begin{cases} H_0 : \text{فرآیند در کنترل است} \\ H_1 : \text{فرآیند خارج از کنترل است} \end{cases}$$

با فرض $X \sim N(\mu_0, \sigma_0)$ اگر μ_0 و σ_0 به $\mu_1 = \mu_0 + k\sigma_0$ ($\sigma_1^2 = k\sigma_0^2$) تغییر پیدا کند، آنگاه نسبت محصولات غیر منطبق و میانگین و واریانس درجه عدم انطباق از مقادیر $p_0, \mu_{\widetilde{N}_0}, \sigma_{\widetilde{N}_0}^2$ به مقادیر $p_1, \mu_{\widetilde{N}_1}, \sigma_{\widetilde{N}_1}^2$ تغییر پیدا می‌کند.

$$p_0 = 1 - p(L \leq X \leq U | H_0), \mu_{\widetilde{N}_0} = E(\widetilde{N}(X) | H_0), \sigma_{\widetilde{N}_0}^2 = Var(\widetilde{N}(X) | H_0)$$

$$p_1 = 1 - p(L \leq X \leq U | H_1), \mu_{\widetilde{N}_1} = E(\widetilde{N}(X) | H_1), \sigma_{\widetilde{N}_1}^2 = Var(\widetilde{N}(X) | H_1)$$

$$\beta_{np} = P\{np_0 - 3n\sqrt{\frac{p_0q_0}{n}} < n\widehat{p} < np_0 + 3n\sqrt{\frac{p_0q_0}{n}} | H_1\}$$

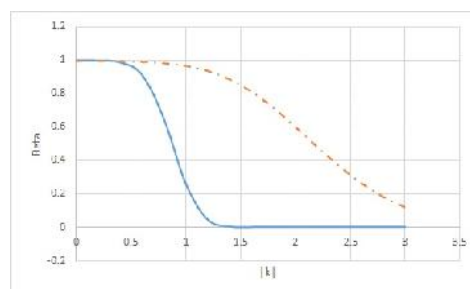
$$\beta_{\widetilde{np}} = P\{n(\mu_{\widetilde{N}_0} - 3\frac{\sigma_{\widetilde{N}_0}^2}{\sqrt{n}}) < n\widetilde{N} < n(\mu_{\widetilde{N}_0} + 3\frac{\sigma_{\widetilde{N}_0}^2}{\sqrt{n}})\}$$

با توجه به میزان تغییرات مختلف در میانگین و واریانس، ابتدا مقدار خطا را در جدول ۲ با توجه به فرمول‌های ذکر شده، محاسبه کردیم. با توجه به شکل ۴ توانایی دو روش را در تشخیص تغییرات در فرآیند با توجه به k های مختلف نشان داده شده است.

جدول ۲: خطای نوع دوم برای $|k|$ مختلف

$ k $	β_{np}	$\beta_{\widetilde{np}}$	$ k $	β_p	$\beta_{\widetilde{p}}$
0	0.9973	0.9976	1.6	0	0.8121
0.2	0.9963	0.9972	1.8	0	0.7147
0.4	0.9850	0.9958	2	0	0.6008
0.6	0.9090	0.9921	2.2	0	0.4805
0.8	0.6434	0.9840	2.4	0	0.3652
1	0.2562	0.9675	2.6	0	0.2639
1.2	0.0441	0.9369	2.8	0	0.1815
1.4	0.0028	0.8864	3	0	0.1192

همانطور که در شکل ۴ مشخص است، میزان خطا در نمودار کنترل np به نسبت بیشتر از میزان خطا نمودار کنترل $\widetilde{np}$ می‌باشد.



شکل ۴: منحنی OC برای نمودار np (—) و نمودار $\bar{np}$ (- -) زمانی که میانگین تغییر می‌کند.

۵. نتیجه گیری

یکی از فرضیه‌های اصلی برای رسم نمودارهای کنترل کیفیت غیرفازی، هم ارزش و یکسان در نظر گرفتن تمامی معایب است. ولی در کنترل کیفیت فازی درجات بی‌کیفیتی/کیفیت برای اقلام معیوب/سالم یکسان در نظر گرفته نمی‌شود. با استفاده از نمودار np دو معیار OC نمودار جدیدی که در این مقاله معرفی شد خیلی کاراتر از نمودار np در حالت کلاسیک است و خطای نوع دوم نمودار $\bar{np}$ کمتر از نمودار کنترل در حالت کلاسیک است.

مراجع

۱. پرچی، ع. (۱۳۹۵). نمودارهای کنترل شوهارت بر اساس کیفیت فازی، اندیشه آماری، ارسال شده برای داوری
۲. مونگمری، د. س. (۱۳۸۹). کنترل کیفیت آماری، ترجمه‌ی رسول نورالسنا، انتشارات دانشگاه علم و صنعت.
3. Amirzadeh, V., Mashinchi, M., & Parchami, A. (2009). *Construction of p-charts using degree of nonconformity*. Information Sciences, 179(1), 150-160.
4. Faraz, A., Baradaran Kazemzadeh, R., Bameni Moghadam, M., & Bazdar, A. (2010). *Constructing a fuzzy Shewhart control chart for variables when uncertainty and randomness are combined*. Quality & Quantity, 44(5), 905-914.
5. Gülbay, M., & Kahraman, C. (2007). *An alternative approach to fuzzy control charts: Direct fuzzy approach*. Information Sciences, 177(6), 1463-1480.
6. Gülbay, M., Kahraman, C., & Ruan, D. (2004). *-cut fuzzy control charts for linguistic data*. Intelligent Systems, 19(12), 1173-1196.
7. Kahraman, A., Tolga, E., & Ulukan, Z. (1995). *Using triangular fuzzy numbers in the tests of control charts for unnatural patterns*. Emerging Technologies and Factory Automation, 3, 291-298.
8. Raz, T., & Wang, J. H. (1990). *Probabilistic and membership approaches in the construction of control charts for linguistic data*. Production Planning & Control, 1(3), 147-157.
9. Wang, j. h., & Raz, T. (1990). *on the construction of control charts using linguistic variables*. production research, 21(3), 477-487.
10. Zadeh, L. A. (1965). *Fuzzy sets*. information and control, 8(3), 338-353.



نمودارهای کنترل کیفیت فازی برای توزیع چوله نرمال

فرزانه هاشمی^{۱*}، ماشاالله ماشینچی^۲، احد جمالیزاده^۳، و مهسا صعصعی^۴

^۱گروه آمار، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان
farzane.hashemi1367@yahoo.com

^۲گروه آمار، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان
mashinchi@uk.ac.ir

^۳گروه آمار، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان
a.jamalizadeh@uk.ac.ir

^۴گروه آمار، دانشکده ریاضی و کامپیوتر، دانشگاه شهید باهنر کرمان
mahsasasaei@gmail.com

چکیده. یکی از ابزارهای مهم کنترل فرایند آماری، نمودارهای کنترل می باشند، نمودارهای کنترلی، طوری طراحی شده اند که تغییرات در میانگین و واریانس یک فرایند تولید را کشف کنند.

نظریه ی مجموعه های فازی، یک روش مفید برای مدلسازی داده های زبانی ارائه می دهد. نمودارهای کنترل فازی نیز، هنگامی که داده های آماری، زبانی هستند مورد استفاده قرار می گیرند. برای طراحی نمودارهای کنترل فازی، می توان داده های زبانی را از طریق روش های غیرفازی سازی به عدد تبدیل نموده و سپس این نمودارها را مورد بررسی قرار داد. در این مقاله نمودارهای کنترل فازی برای داده های چوله نرمال مورد بحث و بررسی قرار می گیرند.

۱. پیش گفتار

کنترل فرایند آماری برخی از روش ها است، که هدفشان تخمین، نظارت و کاهش احتمالی تغییرپذیری در فرایند تولید صنعتی می باشد و نقش مهمی در اطمینان از اینکه فرایند تحت کنترل آماری است، بازی می کند. نمودارهای کنترلی، کاربرد گسترده ای در صنعت، بخصوص در فرایندهای تولیدی دارند. این نمودارهای کنترلی طوری طراحی می شوند تا تغییرات در میانگین و واریانس یک فرایند تولید را کشف کرده و شرایط غیر عادی را نشان دهند. به طور کلی نمودارهای کنترل، فرایند را در صورتی که تحت کنترل یا خارج از کنترل باشد، بررسی می کنند. اگر فرایند تحت کنترل باشد آنگاه تخمین و پیش بینی قابلیت فرایند انجام پذیر است [۸]. با توجه به اینکه نظریه ی مجموعه فازی، قابلیت نشان دادن داده های زبانی را دارد لذا که داده هایی در عباراتی نظیر "خیلی خوب"، "خوب"، "متوسط"، "بد"، "خیلی بد" بیان می شوند، این نظریه ابزار مناسبی برای به کار بردن این داده ها می باشد. بررسی های زیادی انجام شده تا روش های آماری و نظریه ی مجموعه های فازی با همدیگر ترکیب شوند. زاده، طرح کلی تئوری

واژگان کلیدی. چوله نرمال، حدود کنترل فازی، اعداد فازی مثلثی.
* سخنران

تعمیم یافته عدم قطعیت را که تغییر مهم در تبادل بین عدم قطعیت و اطلاعات را نشان می‌دهد بیان می‌کند [۵].

مقالات زیادی راجع به نمودارهای کنترل فازی و کاربردهایشان وجود دارد. وانگ و راز [۱۱] و [۱۲]، حدود کنترل را برای زمانی که مقادیر کلاسیک به مقادیر فازی تغییر می‌کنند پیشنهاد کردند. کاناگوا [۷] یک نمودار کنترل بر اساس تابع چگالی احتمال برای داده‌های زبانی معرفی کرد. رولند و وانگ [۱۰] روش‌های کنترل فرایند آماری (SPC)^۱ فازی بر اساس کاربرد منطق فازی در قواعد SPC معرفی کردند.

هنگامی که داده‌های آماری، زبانی و یا در دسترس نیستند و یا اطلاعات در دسترس، ناقص هستند، نمودارهای کنترل فازی به کار برده می‌شوند [۶].

معمولا در عمل وقتی که مشخصه کیفی مورد مطالعه بصورت متغیر باشد، هم میانگین و هم واریانس آن را کنترل می‌کنیم. در این قسمت نمودارهای کنترل $\bar{X} - R$ برای وقتی که مشخصه‌های کیفیت بصورت اندازه‌های فازی تعریف می‌شوند طراحی می‌شوند.

در نمودارهای کنترل متغیر، برخی اندازه‌های گرایش مرکزی در آمار توصیفی، نظیر مد فازی، میان دامنه فازی سطح α ، میانه فازی و میانگین فازی مورد استفاده قرار می‌گیرند. این اندازه‌ها می‌توانند برای تبدیل مجموعه‌های فازی به مجموعه‌های عددی استفاده شوند. برای انتخاب یکی از این روش‌های غیر فازی سازی پایه تئوری وجود ندارد. این انتخاب اساسا مبنی بر سهولت محاسبه یا اولویت محاسبه کننده می‌باشد. در این مقاله ابتدا در بخش ۲ توزیع چوله نرمال همراه با حدود کنترل آن را معرفی می‌کنیم. در بخش ۳ نمودارهای کنترل فازی $\bar{X} - \bar{R}$ برای اعداد فازی مثلثی^۲ (TFN) در حالت توزیع چوله نرمال بیان می‌کنیم. در پایان نیز یک مثال کاربردی از آن‌ها را ارائه می‌دهیم.

۲. توزیع چوله نرمال و حدود کنترل آن

۱.۲. توزیع چوله نرمال. از گذشته‌های دور، در مباحث آماری خانواده‌های پارامتری و مطالعاتی جزئیات آنها موضوع مورد بحث صاحب‌نظران بوده و در چند سال اخیر، این مبحث با موفقیت چشمگیری همراه بوده است. قسمت قابل توجهی از این مطالعات، درباره خانواده توزیع چوله نرمال می‌باشد. هنگامی که اصطلاح چوله نرمال را به کار می‌بریم منظور یک کلاس از توزیع‌های پارامتری احتمال است که توزیع نرمال را با اضافه کردن پارامتری به نام پارامتر شکل^۳ که بیان کننده چولگی است، به کلاس دیگری از توزیع‌ها که توزیع نرمال یک حالت خاص از آن می‌باشد تعمیم و از نرمال بودن خارج می‌سازد. از لحاظ عملی توزیع چوله نرمال به عنوان تعمیمی از توزیع نرمال در موقعیت‌هایی کاربرد دارد که داده‌ها دارای مقداری چولگی یا کجی باشند و توزیع نرمال بخوبی روی داده‌ها برازش نشود. از لحاظ تئوری توزیع چوله نرمال علاوه بر حفظ کردن بسیاری از ویژگی‌های توزیع نرمال، خود دارای خصوصیات مهم و جالب توجهی است که این توزیع را به

^۱Statistical Process Control

^۲Triangular fuzzy number

^۳shape parameter

یک توزیع پر کاربرد تبدیل کرده است. توزیع چوله نرمال ابتدا توسط آزالینی^۴ [۱] معرفی شد. در ادامه آزالینی و دالا واله^۵ [۲] حالت چند متغیره این توزیع را ارایه کردند. همچنین آزالینی و کاپیتانو^۶ [۳] خواص مهم و احتمالی این توزیع را بدست آوردند.

تعریف ۱.۲. [۱] گوییم متغیر تصادفی Z_λ دارای توزیع چوله نرمال است و با نماد $Z_\lambda \sim SN(\xi, \sigma, \lambda)$ نشان می‌دهیم، اگر شکل تابع چگالی آن به صورت زیر باشد:

$$f_X(x) = \frac{2}{\sigma} \phi\left(\frac{x-\xi}{\sigma}\right) \Phi\left(\lambda \frac{x-\xi}{\sigma}\right), \quad \xi \in \mathbb{R}, \sigma \in \mathbb{R}^+, \lambda \in \mathbb{R},$$

به طوریکه $\phi(\cdot)$ و $\Phi(\cdot)$ به ترتیب، تابع چگالی و تابع توزیع نرمال استاندارد می‌باشند.

قضیه ۲.۲. [۱] اگر متغیر تصادفی $Z_\lambda \sim SN(\lambda)$ باشد آنگاه

$$\begin{aligned} \mu_x &= \xi + a_1 \rho \sigma, \\ \sigma_x^2 &= \sigma^2 (1 - a_1^2 \rho^2), \end{aligned}$$

به طوریکه $a_1 = \frac{\lambda}{\sqrt{1+\lambda^2}}$ ؛ $\rho = \frac{\lambda}{\sqrt{1+\lambda^2}}$ ، ξ ، σ و λ به ترتیب پارامترهای مکان، مقیاس و چولگی می‌باشند.

لم ۳.۲. [۴] متغیر تصادفی $Z_1, Z_2, \dots, Z_n \sim SN(\lambda)$ باشد، آنگاه

$$f_{\bar{Z}}(z; n, \lambda) = 2^n \phi\left(z; 0, \frac{1}{n}\right) \Phi_n\left(\frac{n\lambda z}{\sqrt{1+\lambda^2}}, 0, I_n - \frac{\lambda^2}{(1+\lambda^2)n} I_n\right),$$

$$-\infty < z < \infty$$

بطوریکه $\Phi_n(\cdot)$ تابع توزیع نرمال n متغیره می‌باشد.

۱.۱.۲. نمودارهای کنترل برای داده‌های چوله نرمال. به طور معمول وقتی مشخصه کیفی مورد مطالعه به صورت متغیر باشد هم میانگین و هم واریانس آن را کنترل می‌کنیم. میانگین فرایند اغلب به وسیله نمودار کنترل میانگین یا نمودار $\bar{X}$ کنترل می‌شود. تغییرپذیری فرایند را می‌توان به وسیله نمودار کنترل برای انحراف معیار یا نمودار کنترل برای دامنه که به ترتیب نمودار S و نمودار R نامیده می‌شوند کنترل کرد. نمودار کنترل R به علت سادگی در عمل بیشتر مورد استفاده قرار می‌گیرد. به طور معمول برای هر مشخصه کیفی یک نمودار کنترل $\bar{X}$ و R مجزا استفاده می‌شود. نمودارهای کنترل $\bar{X}$ و R (یا S) از مهم‌ترین ابزار کنترل فرایند آماری هستند. باید توجه داشت که کنترل هر دو مقدار میانگین و پراکندگی فرایند از اهمیت ویژه‌ای برخوردار است. اولین بار تسای^۷ در سال ۲۰۰۷ نمودارهای کنترل را برای نظارت بر میانگین فرایند، وقتی که داده‌ها دارای توزیع چوله نرمال هستند، پیشنهاد کرد. ولی تسای توزیع میانگین نمونه

^۴Azzalini

^۵Dalla Valla

^۶Capitiano

^۷Tsai

را به طور تقریبی محاسبه کرده بود، و این انگیزه‌ای شد که چانگ^۱ در سال ۲۰۱۴ توزیع دقیق میانگین نمونه را بدست آورد و نمودارهای کنترل را برای داده‌های چوله نرمال توسعه دهد [۴].

قضیه ۴.۲. [۴] حدود نمودار کنترل $\bar{X}$ و R برای داده‌های چوله نرمال به صورت زیر می‌باشند:

$$\begin{aligned} UCL_{\bar{X}} &= \mu_{x_0} + \left(\gamma_1 \bar{z}_{\lambda, n, 1 - \frac{\alpha}{2}} - \gamma_2 \right) \sigma_{x_0}, \\ LCL_{\bar{X}} &= \mu_{x_0} + \left(\gamma_1 \bar{z}_{\lambda, n, \frac{\alpha}{2}} - \gamma_2 \right) \sigma_{x_0}. \end{aligned} \quad (۱.۲)$$

به طوریکه μ_{x_0} و σ_{x_0} میانگین و واریانس فرایند تحت کنترل هستند. و $\gamma_1 = (1 - a_1^2 \rho^2)^{-\frac{1}{2}}$ و $\gamma_2 = \gamma_1 a_1 \rho$ ، $\bar{z}_{\lambda, n, \frac{\alpha}{2}}$ و $\bar{z}_{\lambda, n, 1 - \frac{\alpha}{2}}$ چنک $100(1 - \frac{\alpha}{2})$ توزیع چوله نرمال استاندارد می‌باشند.

$$\begin{aligned} UCL_R &= \sigma_{x_0} \gamma_1 r_{\lambda, n, 1 - \frac{\alpha}{2}}, \\ LCL_R &= \sigma_{x_0} \gamma_1 r_{\lambda, n, \frac{\alpha}{2}}. \end{aligned} \quad (۲.۲)$$

برای محاسبه‌ی $r_{\lambda, n, \alpha}$ به توزیع R_Z نیاز داریم.

لم ۵.۲. [۴] فرض کنید $Z_{(1)}$ و $Z_{(n)}$ اولین و آخرین آماری ترتیبی از نمونه‌ی $Z_1, \dots, Z_n$ باشند. تابع چگالی توام $Z_{(1)}$ و $Z_{(n)}$ به صورت زیر است:

$$f_{Z_{(n)}, Z_{(1)}}(x, y) = n(n-1) [F_Z(x) - F_Z(y)]^{n-2} f_Z(x) f_Z(y), \quad x > y,$$

به طوریکه $f_Z(x)$ و $F_Z(x)$ تابع چگالی و تابع توزیع چوله نرمال استاندارد می‌باشند. تابع چگالی R_Z به صورت زیر بدست می‌آید:

$$f_{R_Z}(r) = \int_{-\infty}^{\infty} n(n-1) [F_Z(r+y) - F_Z(y)]^{n-2} f_Z(r+y) f_Z(y) dy, \quad r > 0$$

با توجه به ۵.۲، $r_{\lambda, n, \alpha}$ به روش عددی با استفاده از f_{R_Z} محاسبه می‌شود.

لم ۶.۲. [۴] در صورتی که پارامترها نامعلوم باشند حدود کنترل بصورت زیر می‌باشند:

$$\begin{aligned} UCL_{\bar{X}} &= \bar{\bar{X}} + (\gamma_1 \bar{z}_{\lambda, n, 0.99865} - \gamma_2) \frac{\bar{R}}{\gamma_1 d_2^*} = \bar{\bar{X}} + A_2 \bar{R} \\ LCL_{\bar{X}} &= \bar{\bar{X}} + (\gamma_1 \bar{z}_{\lambda, n, 0.00135} - \gamma_2) \frac{\bar{R}}{\gamma_1 d_2^*} = \bar{\bar{X}} + A_1 \bar{R} \end{aligned} \quad (۳.۲)$$

و

$$\begin{aligned} UCL_R &= \gamma_1 r_{\lambda, n, 0.99865} \frac{\bar{R}}{\gamma_1 d_2^*} = \bar{R} D_4 \\ LCL_R &= \gamma_1 r_{\lambda, n, 0.00135} \frac{\bar{R}}{\gamma_1 d_2^*} = \bar{R} D_3 \end{aligned} \quad (۴.۲)$$

^۱Chung-I Li

به طوریکه

$$\begin{aligned}d_2^* &= E\left(\frac{R}{\sigma_{x_0}\gamma_1}\right), \\ \bar{X} &= \frac{\sum_{i=1}^n \sum_{j=1}^m X_{ij}}{nm}, \\ A_2 &= (\gamma_1 \bar{z}_{\lambda,n,0.99865} - \gamma_2) \frac{1}{\gamma_1 d_2^*}, \\ A_1 &= (\gamma_1 \bar{z}_{\lambda,n,0.00135} - \gamma_2) \frac{1}{\gamma_1 d_2^*}, \\ D_3 &= \gamma_1 r_{\lambda,n,0.99865} \frac{1}{\gamma_1 d_2^*}\end{aligned}$$

و

$$D_4 = \gamma_1 r_{\lambda,n,0.99865} \frac{1}{\gamma_1 d_2^*}$$

می‌باشند.

۳. نمودارهای کنترل فازی برای اعداد فازی مثلثی

تعریف ۱.۳. [۹] فرض کنید، مشخصه‌ی کیفیت بصورت ”تقریباً x ” تعریف شده باشد. در این صورت، این مقادیر زبانی را می‌توان به صورت اعداد فازی مثلثی، به شکل $\tilde{x} = (x_1, x_2, x_3)$ تعریف کرد.

با توجه به تعریف ۱.۳ برای نمونه تصادفی از اعداد فازی $\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_n$ میانگین اعداد فازی مثلثی به صورت زیر محاسبه می‌شود.

$$\tilde{\bar{x}} = TFN\left(\frac{\sum_{j=1}^n x_{1j}}{n}, \frac{\sum_{j=1}^n x_{2j}}{n}, \frac{\sum_{j=1}^n x_{3j}}{n}\right) = TFN(\bar{x}_1, \bar{x}_2, \bar{x}_3). \quad (1.3)$$

به همین ترتیب برای نمونه تصادفی از اعداد فازی $\tilde{x}_1, \tilde{x}_2, \dots, \tilde{x}_n$ دامنه نمونه اعداد فازی مثلثی به صورت زیر بدست می‌آید.

$$\begin{aligned}\tilde{R} = TFN\left((\max_j x_{1j} - \min_j x_{3j}), (\max_j x_{2j} - \min_j x_{2j})\right. \\ \left., (\max_j x_{3j} - \min_j x_{1j})\right) = TFN(R_1, R_2, R_3). \quad (2.3)\end{aligned}$$

به طوریکه $(\max_j x_{1j}, \max_j x_{2j}, \max_j x_{3j})$ و $(\min_j x_{1j}, \min_j x_{2j}, \min_j x_{3j})$ ، به ترتیب مقادیر ماکسیمم و مینیمم اعداد فازی مثلثی را نمایش می‌دهند.

تعریف ۲.۳. [۹] حال فرض می‌کنیم m نمونه تصادفی هر یک به حجم n از فرایند تولید انتخاب شود، در این صورت، میانگین میانگین‌های فازی ($\tilde{X}$) و میانگین دامنه‌ی نمونه فازی ($\tilde{R}$)، به شکل زیر قابل محاسبه هستند

$$\begin{aligned}\tilde{\bar{x}} &= TFN\left(\frac{\sum_{j=1}^m \bar{x}_{1j}}{m}, \frac{\sum_{j=1}^m \bar{x}_{2j}}{m}, \frac{\sum_{j=1}^m \bar{x}_{3j}}{m}\right) = TFN(\bar{x}_1, \bar{x}_2, \bar{x}_3) \\ &= TFN(\hat{\mu}_1, \hat{\mu}_2, \hat{\mu}_3).\end{aligned}\quad (۳.۳)$$

$$\tilde{\bar{R}} = \left(\frac{\sum_{j=1}^m R_{1j}}{m}, \frac{\sum_{j=1}^m R_{2j}}{m}, \frac{\sum_{j=1}^m R_{3j}}{m}\right) = TFN(\bar{R}_1, \bar{R}_2, \bar{R}_3). \quad (۴.۳)$$

لم ۳.۳. [۹] حدود کنترل فازی، برای نمودار کنترل $\tilde{X} - \tilde{R}$ ، برای داده‌های چوله نرمال به شکل زیر قابل محاسبه است:

$$\begin{aligned}U\tilde{C}L_{\tilde{X}} &= \tilde{\bar{X}} + A_2\tilde{\bar{R}} = (\bar{x}_1 + A_2\bar{R}_1, \bar{x}_2 + A_2\bar{R}_2, \bar{x}_3 + A_2\bar{R}_3) \\ &= TFN(UCL_1\bar{x}, UCL_2\bar{x}, UCL_3\bar{x}) \\ \tilde{C}L_{\tilde{X}} &= \tilde{\bar{X}} = (\bar{x}_1, \bar{x}_2, \bar{x}_3) \\ &= TFN(CL_1\bar{x}, CL_2\bar{x}, CL_3\bar{x})\end{aligned}\quad (۵.۳)$$

$$\begin{aligned}L\tilde{C}L_{\tilde{X}} &= \tilde{\bar{X}} + A_1\tilde{\bar{R}} = (\bar{x}_1 + A_1\bar{R}_1, \bar{x}_2 + A_1\bar{R}_2, \bar{x}_3 + A_1\bar{R}_3) \\ &= TFN(LCL_1\bar{x}, LCL_2\bar{x}, LCL_3\bar{x})\end{aligned}$$

مقدار A_2 و A_1 از رابطه‌ی (۳.۲) محاسبه می‌شود، این مقدار برای اندازه نمونه‌های مختلف در [۴] قابل دستیابی است.

حدود کنترل فازی، برای نمودار کنترل $\tilde{R}$ به شکل زیر قابل محاسبه است:

$$\begin{aligned}U\tilde{C}L_R &= \tilde{\bar{R}}D_4 = TFN(R_1D_4, R_2D_4, R_3D_4) \\ &= TFN(UCL_1R, UCL_2R, UCL_3R)\end{aligned}\quad (۶.۳)$$

$$\tilde{C}L_R = \tilde{\bar{R}} = TFN(R_1, R_2, R_3) = TFN(CL_1R, CL_2R, CL_3R)$$

$$\begin{aligned}L\tilde{C}L_R &= \tilde{\bar{R}}D_3 = TFN(R_1D_3, R_2D_3, R_3D_3) \\ &= TFN(LCL_1R, LCL_2R, LCL_3R)\end{aligned}$$

مقادیر D_3 و D_4 از رابطه‌ی (۴.۲)، محاسبه می‌شوند، این مقادیر برای اندازه نمونه‌های مختلف در [۴] قابل دستیابی هستند.

۴. مثال عددی

فرض کنید در یک فرایند تولید در ساعت‌های مختلف نمونه‌هایی به حجم ۴ گرفته شده است. این کار ۷ مرتبه تکرار شده است، و این داده‌ها با توجه به [۹] در جدول ۱ داده شده‌اند. حال در نظر داریم با استفاده از حدود کنترل فازی برای توزیع چوله نرمال بررسی کنیم که آیا فرایند با توجه

به مقادیر مشاهدات، که اعداد فازی مثلثی هستند، تحت کنترل است یا خیر برای این کار ابتدا میانگین میانگین‌های فازی $(\tilde{\bar{X}})$ و میانگین دامنه‌ی نمونه فازی $(\tilde{\bar{R}})$ را طبق روابط (۳.۳) و (۴.۳) به دست می‌آوریم. مقدار $m = 7$ می‌باشد:

$$\begin{aligned}\tilde{\bar{x}} &= TFN\left(\frac{\sum_{j=1}^m \bar{x}_{1j}}{m}, \frac{\sum_{j=1}^m \bar{x}_{2j}}{m}, \frac{\sum_{j=1}^m \bar{x}_{3j}}{m}\right) = TFN(\bar{\bar{x}}_1, \bar{\bar{x}}_2, \bar{\bar{x}}_3) \\ &= TFN(\hat{\mu}_1, \hat{\mu}_2, \hat{\mu}_3) = TFN(74.00, 74.012, 74.018).\end{aligned}$$

$$\begin{aligned}\tilde{\bar{R}} &= TFN\left(\frac{\sum_{j=1}^m R_{1j}}{m}, \frac{\sum_{j=1}^m R_{2j}}{m}, \frac{\sum_{j=1}^m R_{3j}}{m}\right) = TFN(\bar{R}_1, \bar{R}_2, \bar{R}_3) \\ &= TFN(0.027, 0.041, 0.043).\end{aligned}$$

برای اندازه نمونه‌های $n = 4$ با استفاده از [۴] مقدار $D_4 = 2.73$ و $D_3 = 0.1$ به دست می‌آید. بنابراین با استفاده از روابط (۶.۳) حدود کنترل فازی برای توزیع چوله نرمال برای نمودار $\tilde{R}$ به شکل زیر حاصل می‌شود:

$$\begin{aligned}U\tilde{C}L_R &= \tilde{\bar{R}}D_4 = TFN(R_1D_4, R_2D_4, R_3D_4) \\ &= TFN(UCL_1R, UCL_2R, UCL_3R) \\ &= TFN(0.0737, 0.11310, 0.1178) \\ \tilde{C}L_R &= \tilde{\bar{R}} = TFN(R_1, R_2, R_3) = TFN(CL_1R, CL_2R, CL_3R) \\ &= TFN(0.0270, 0.0414, 0.0431) \\ L\tilde{C}L_R &= \tilde{\bar{R}}D_3 = TFN(R_1D_3, R_2D_3, R_3D_3) \\ &= TFN(LCL_1R, LCL_2R, LCL_3R) \\ &= TFN(0.0027, 0.0041, 0.0043)\end{aligned}$$

برای اندازه نمونه‌های $n = 4$ با استفاده از [۴] مقدار $D_2^U = 0.8$ و $A_2^L = 0.6$ به دست می‌آید. بنابراین با استفاده از روابط (۵.۳) حدود کنترل فازی برای توزیع چوله نرمال برای نمودار

جدول ۱: نمونه‌ی ۷ تایی با اندازه‌ی ۴

$\tilde{X}_1 = (X_{11}, X_{21}, X_{31})$	$\tilde{X}_2 = (X_{12}, X_{22}, X_{32})$
(73.980, 73.989, 73.999)	(73.999, 74.001, 74.002)
(74.001, 74.002, 74.003)	(73.970, 73.980, 73.990)
(73.983, 73.985, 73.986)	(73.989, 73.990, 73.991)
(74.001, 74.003, 74.005)	(74.002, 74.004, 74.007)
(73.990, 74.001, 74.010)	(74.041, 74.042, 74.050)
(74.009, 74.010, 74.012)	(73.999, 74.007, 74.009)
(73.970, 73.980, 74.000)	(74.020, 74.029, 74.040)
$\tilde{X}_3 = (X_{13}, X_{23}, X_{33})$	$\tilde{X}_4 = (X_{14}, X_{24}, X_{34})$
(74.019, 74.029, 74.039)	(74.020, 74.030, 74.040)
(73.999, 74.003, 74.005)	(74.050, 74.059, 74.069)
(74.002, 74.004, 74.005)	(74.006, 74.007, 74.008)
(74.003, 74.004, 74.005)	(73.999, 74.000, 74.002)
(74.047, 74.053, 74.061)	(74.025, 74.035, 74.045)
(73.999, 74.001, 74.002)	(74.005, 74.006, 74.007)
(74.029, 74.031, 74.033)	(74.064, 74.065, 74.066)

$\tilde{\tilde{X}}$ به شکل زیر حاصل می‌شود:

$$\begin{aligned} U\tilde{C}L_{\tilde{X}} &= \tilde{\tilde{X}} + A_2\tilde{\tilde{R}} = TFN(\bar{x}_1 + A_2\bar{R}_1, \bar{x}_2 + A_2\bar{R}_2, \bar{x}_3 + A_2\bar{R}_3) \\ &= TFN(UCL_1\bar{x}, UCL_2\bar{x}, UCL_3\bar{x}) \\ &= TFN(74.0214, 74.0459, 74.0520) \end{aligned}$$

$$\begin{aligned} \tilde{C}L_{\tilde{X}} &= \tilde{\tilde{X}} = TFN(\bar{x}_1, \bar{x}_2, \bar{x}_3) = TFN(CL_1\bar{x}, CL_2\bar{x}, CL_3\bar{x}) \\ &= TFN(73.9998, 74.0128, 74.0175) \end{aligned}$$

$$\begin{aligned} L\tilde{C}L_{\tilde{X}} &= \tilde{\tilde{X}} + A_1\tilde{\tilde{R}} = TFN(\bar{x}_1 + A_1\bar{R}_1, \bar{x}_2 + A_1\bar{R}_2, \bar{x}_3 + A_1\bar{R}_3) \\ &= TFN(LCL_1\bar{x}, LCL_2\bar{x}, LCL_3\bar{x}) \\ &= TFN(73.9836, 73.9879, 73.9916) \end{aligned}$$

پس از محاسبه‌ی حدود کنترل طبق روش قوانین فازی برای اعداد فازی مثلثی [۹] عمل می‌کنیم و در مورد فرایند تصمیم می‌گیریم. که نتایج حاصل در جدول ۲ داده شده است. همچنین تمامی محاسبات و نتایج حاصل در این بخش، از طریق برنامه نویسی در نرم افزار R انجام شده است.

مراجع

1. Azzalini, A., (1985). *A class of distributions which includes the normal ones*. Scand. J. Statist. 12, 171-178.
2. Azzalini, A. and Dalla Valle, A. (1996). *The multivariate skew-normal distribution*. Biometrika 83, 715-726.
3. Azzalini, A. and Capitanio, A. (1999). *Statistical applications of the multivariate skew normal distribution*. J. Roy. Statist. Soc. Ser. B 61, 579-602.

جدول ۲: نتایج کنترل فرآیند در ۷ نمونه

نمونه	قوانین فازی برای نمودار $\bar{X}$	قوانین فازی برای نمودار $\bar{R}$	کنترل فرآیند
۱	تحت کنترل	تحت کنترل	تحت کنترل
۲	تحت کنترل	تحت کنترل	تحت کنترل
۳	تحت کنترل	تحت کنترل	تحت کنترل
۴	تقریباً تحت کنترل	تحت کنترل	تقریباً تحت کنترل
۵	تحت کنترل	تحت کنترل	تحت کنترل
۶	تحت کنترل	تحت کنترل	تحت کنترل
۷	تحت کنترل	تحت کنترل	تحت کنترل

4. Chung-I Li, Nan-Cheng Su, Pei-Fang Su, and Yu Shyr *the design of $\bar{X}$ and R control charts for Skew Normal distributed data*, 43, (2012), 4908-4924.
5. Dubois D. and Prade H. (1980), *Fuzzy Sets and Systems*, Theory and applications, Academic.
6. Grzegorzewsk. P (1997), *Control charts for fuzzy data*, In Proceedings of the 5th European congress. on intelligent techniques and soft computing EUFIT97, Aachen, 1326-1330.
7. Kanagawa.A, Tamaki. F, Ohta.H, *Control charts for process average and variability based on linguistic data*, Intelligent Journal of Production Research, 31(4) (1993), 913-922.
8. Kaya. I, Kahraman.C, (2011), *Process capability analyses based on fuzzy measurements and fuzzy control charts*, Expert Systems with Applications, 38 (2011), 3172-3184.
9. Khademi .M, Amirzadeh.V, *fuzzy rules for fuzzy $\bar{X}$ and R control charts*, Iranian Journal of Fuzzy Systems, Vol. 11, No. 5, (2014) pp. 55-66.
10. Rowlands.H, Wang.L.R, *An approach of fuzzy logic evaluation and control in SPC*, Quality Reliability Engineering Intelligent, 16 (2000), 91-98.
11. Wang.J. H, Raz.T, *Applying fuzzy set theory in the development of quality control chart*, Proceeding of the International of Production Research, 28 (1988), 30-35.
12. Wang.J. H, Raz.T, *on the construction of control charts using linguistic variables*, International Journal of Production Research, 28 (1990), 477-487.



رگرسیون نیمه پارامتری فازی بر اساس خوشه‌بندی فازی

معصومه اسداللهی^۱ * و محمدقاسم اکبری^۱

^۱دانشکده علوم ریاضی و آمار، دانشگاه بیرجند

m.asadollahi@birjand.ac.ir

چکیده. در این مقاله، به تخمین پارامترهای مدل رگرسیون نیمه پارامتری با استفاده از متری که بر اساس تفاضل α -شک‌ها است، پرداخته می‌شود. در این رویکرد، که در آن از کمترین قدرمطلق خطا با استفاده از روش‌های خوشه‌بندی فازی استفاده می‌شود یک روش رگرسیون فازی معرفی می‌شود. در انتها نتایج حاصل را با روش یانگ و کو بر اساس معیارهای نیکویی برازش پیشنهادی، مقایسه می‌کنیم.

۱. مقدمه

برای دسته‌بندی افراد، اشیاء و داده‌ها در گروه‌های (خوشه‌های) مختلف از روش‌های خوشه‌بندی استفاده می‌شود که در آن هر خوشه شامل تعدادی از افراد یا اشیاء است که دارای خصوصیات مشابهی باشند. خوشه‌بندی از جمله روش‌های پرکاربرد در تجزیه و تحلیل داده‌های چند متغیره است. در روش‌های خوشه‌بندی کلاسیک، هر داده به یک و فقط یک خوشه نسبت داده می‌شود، در حالی‌که در خوشه‌بندی فازی، یک تفکیک با افراز فازی صورت می‌گیرد، به این معنی که هر داده با یک درجه عضویت به یک خوشه متعلق است.

واژگان کلیدی. عدد فازی، آلفا-شک، خوشه‌بندی فازی، رگرسیون تعمیم‌یافته‌ی یک مرحله‌ای، رگرسیون نیمه

پارامتری .

* سخنران

پس از تعریف مجموعه‌های فازی توسط زاده (۱۹۵۶)، اولین گام در این جهت توسط رسپینی (۱۹۶۹) برداشته شد. یانگ و کو (۱۹۹۶)، یک روش خوشه‌بندی فازی را ارائه کردند. این روش تعمیمی از روش خوشه‌بندی میانگین معمولی، برای حالتی است که داده‌ها به صورت فازی مشاهده شده باشند.

در این مقاله، پس از بیان مفاهیم اولیه، به تخمین پارامترهای مدل رگرسیون نیمه پارامتری با استفاده از متری که بر اساس تفاضل α -شک‌ها است، پرداخته می‌شود. در این رویکرد، که در آن از کمترین قدرمطلق خطا با استفاده از روش‌های خوشه‌بندی فازی استفاده می‌شود یک روش رگرسیون فازی معرفی می‌شود. در انتها نتایج حاصل را با روش یانگ و کو بر اساس معیارهای نیکویی برازش پیشنهادی، مقایسه می‌کنیم.

۲. مفاهیم اولیه

در این بخش به بیان برخی مفاهیم اساسی که در ادامه مقاله مورد نیاز است، می‌پردازیم.

تعریف ۱.۲. مجموعه‌ی فازی $\tilde{A}$ از مجموعه‌ی مرجع X به صورت $\{(x, \tilde{A}(x)) | x \in X\}$ تعریف می‌شود، که $\tilde{A}(x) : X \rightarrow [0, 1]$ تابعی است که برای هر x میزان عضویت آن را در مجموعه‌ی A نشان می‌دهد.

تعریف ۲.۲. مجموعه‌ی فازی $\tilde{A}$ از $\mathbb{R}$ مفروض است. اگر $\tilde{A}$ تک‌نمایی باشد و α -برش‌های آن به ازای هر $\alpha \in [0, 1]$ بازه‌های بسته و کراندار باشند، $\tilde{A}$ را یک عدد فازی گوئیم.

تعریف ۳.۲. اگر عدد فازی $\tilde{A}$ دارای تابع عضویتی به صورت زیر باشد

$$\tilde{A}(x) = \begin{cases} L(\frac{m-x}{l}), & x \leq m, \\ R(\frac{x-m}{r}), & x > m, \end{cases}$$

که در آن L و R توابعی غیر صعودی از $\mathbb{R}^+$ به $[0, 1]$ هستند و $L(0) = R(0) = 1$ ، آن‌گاه $\tilde{A}$ را یک عدد فازی LR نامیده و با نماد $\tilde{A} = (m; l, r)_{LR}$ نمایش می‌دهیم. عدد حقیقی m

مقدار نما (میان)، و اعداد مثبت l و r به ترتیب پهنای چپ و پهنای راست $\tilde{A}$ نامیده می‌شوند. R و L توابع مرجع (شکل) نامیده می‌شود. اگر $\tilde{A}, R(x) = L(x) = \max\{0, 1 - x\}$ را یک عدد فازی مثلثی نامیده و با $\tilde{A} = (m; l, r)_T$ نشان می‌دهیم. در صورتی که $L = R$ و $l = r$ آن‌گاه عدد فازی را متقارن نامیده و با $\tilde{A} = (m; l)_L$ نمایش می‌دهیم. مجموعه‌ی کل اعداد فازی LR را با $F_{LR}(\mathbb{R})$ نشان می‌دهیم.

قضیه ۴.۲. اگر $\tilde{A} = (m; l, r)$ و $\tilde{B} = (n; l', r')$ و $\lambda \in R$ آن‌گاه حساب اعداد فازی را به صورت زیر می‌توان نوشت
(طاهری و کلکین‌نما (۲۰۱۲))

$$\tilde{A} \oplus \tilde{B} = (m; l, r)_{LR} \oplus (n; l', r')_{LR} = (m + n; l + l', r + r')_{LR},$$

$$\tilde{A} \ominus \tilde{B} = (m; l, r)_{LR} \ominus (n; l', r')_{RL} = (m - n; l + r', r + l')_{LR},$$

$$\lambda \otimes \tilde{A} = \begin{cases} (\lambda m; \lambda l, \lambda r)_{LR}, & \lambda > 0, \\ (\lambda m; -\lambda r, -\lambda l)_{RL}, & \lambda < 0. \end{cases}$$

تعریف ۵.۲. لیو (۲۰۱۳) معیاری را برای مقایسه‌ی بین اعداد فازی و اعداد حقیقی معرفی نموده است که از آن به عنوان درجه اعتبار (میزان کوچکی عدد فازی $\tilde{A}$ نسبت به مقدار حقیقی x) یاد می‌شود. این معیار به صورت تابع $C : F(\mathbb{R}) \rightarrow [0, 1]$ تعریف می‌شود، که در آن

$$C(\tilde{A} \leq x) = \frac{\sup_{y \leq x} \tilde{A}(y) + 1 - \sup_{y > x} \tilde{A}(y)}{2}.$$

تعریف ۶.۲. بزرگ‌ترین کران پایین مجموعه‌ی تمام عناصری از تکیه‌گاه مجموعه‌ی فازی $\tilde{A}$ را که تابع درجه اعتبار آن‌ها حداقل به بزرگی α باشد، α -شک $\tilde{A}$ نامیده و با $\tilde{A}_\alpha$ نمایش می‌دهیم، یعنی

$$\tilde{A}_\alpha = \inf\{x \in \text{supp}(\tilde{A}) : C(\tilde{A} \leq x) \geq \alpha\}$$

که $\tilde{A}_\alpha$ برای $\alpha \in (0, 1]$ یک تابع غیر نزولی است.

تعریف ۷.۲. فرض کنید $\tilde{M} = (m; l, r)$ ، $\tilde{N} = (n; \acute{l}, \acute{r})$ دو عدد فازی LR باشند، در این صورت فاصله‌ی بین آن‌ها براساس $-\alpha$ شک را با K نمایش داده و به صورت زیر تعریف می‌کنیم

$$K(\tilde{M}, \tilde{N}) = \int_0^1 |\tilde{M}_\alpha - \tilde{N}_\alpha| d\alpha$$

حال اگر $\tilde{M}$ و $\tilde{N}$ دو عدد فازی مثلی غیر متقارن باشند داریم

$$K(\tilde{M}, \tilde{N}) = \frac{1}{2} \int_0^1 |(m - n) + \alpha(l' - l)| d\alpha + \frac{1}{2} \int_0^1 |(m - n) + \alpha(r - r')| d\alpha$$

تعریف ۸.۲. برای دو عدد فازی $\tilde{M} = (m; l, r)_{LR}$ و $\tilde{N} = (n; \acute{l}, \acute{r})_{LR}$ اگر $l \geq \acute{l}$ و $r \geq \acute{r}$ تفاضل ها کوها را به صورت زیر تعریف می‌شود

$$\tilde{M} \ominus_H \tilde{N} = (m - n; l - \acute{l}, r - \acute{r})_{LR}$$

که در صورت عدم شرایط فوق، این تفاضل وجود نخواهد داشت. (اکبری (۱۳۸۷))

۳. رگرسیون نیمه پارامتری در محیط فازی

مدل رگرسیون خطی نیمه پارامتری در حالت تک متغیره به صورت $Y = X\beta + f(T) + \epsilon$ تعریف می‌شود، در این مدل فرض می‌شود که تابع ناپارامتری $f(\cdot)$ یک تابع هموار و T متغیرهای مرتب شده به طور صعودی در یک بازه‌ی کراندار است و برای خطاهای مستقل از هم، $E(\epsilon|x, t) = 0$ و $Var(\epsilon|x, t) = \sigma_\epsilon^2$ می‌باشند. رابینسن (۱۹۸۸)، یک برآورد برای پارامترهای خطی این مدل بر اساس برآورد تابع غیر خطی با روش هموارساز کرنل پیشنهاد کرد.

در این بخش تعمیمی از روش رابینسن (۱۹۸۸) برای برآورد پارامتر خطی مدل رگرسیون نیمه پارامتری در محیط فازی با پارامترهای خطی غیر فازی و داده‌های متغیر مستقل و متغیر وابسته‌ی فازی ارائه می‌شود.

یک مدل رگرسیون نیمه پارامتری یک متغیره در محیط فازی به صورت زیر تعریف می شود
(اسداللهی (۱۳۹۵))

$$\tilde{y} = a_{\setminus} \otimes \tilde{x} \oplus \tilde{f}(t) \quad (۱.۳)$$

که در آن $\tilde{f}(t_i) = (f_i; l_{f_i}, r_{f_i})_{LR}$ و $\tilde{x}_i = (x_i; l_{x_i}, r_{x_i})_{LR}$ ، $\tilde{y}_i = (y_i; l_{y_i}, r_{y_i})_{LR}$ برای سادگی فرض می کنیم که $t_i = \frac{i}{n}, (i = ۱, ۲, \dots, n)$ به صورت صعودی روی بازه $[۰, ۱]$ توزیع شده اند. با استفاده از روش هموارساز کرنل برآورد $\tilde{f}(t_i)$ را به شکل زیر به دست می آوریم

$$\begin{aligned} \hat{\tilde{f}}(t_i) &= \sum_{j=1}^n w_j(t_i) \otimes \left(\tilde{y}_j \ominus_H (a_{\setminus} \otimes \tilde{x}_j) \right) \\ &= \sum_{j=1}^n w_j(t_i) \otimes \left(\tilde{y}_j \ominus_H \sum_{j=1}^n w_j(t_i) a_{\setminus} \otimes \tilde{x}_j \right). \end{aligned}$$

که در آن $w_j(t_i) = \frac{\frac{1}{nh} K(\frac{t_j - t_i}{h})}{\frac{1}{nh} \sum_{j=1}^n K(\frac{t_j - t_i}{h})}$ وزن ها توسط تابع کرنل بکار رفته $K(\cdot)$ و مقادیر آن ها بر اساس فاصله از نقطه t_i و پهنای باند h تعیین می شود.

با جایگذاری این برآورد در رابطه‌ی (۱.۳) به رابطه‌ی زیر می‌رسیم

$$\begin{aligned}\tilde{y}_i &= (a_1 \otimes \tilde{x}_i) \oplus \sum_{j=1}^n (w_j(t_i) \otimes \tilde{y}_j) \ominus_H \sum_{j=1}^n \left(w_j(t_i) a_1 \otimes \tilde{x}_j \right) \\ &= \sum_{j=1}^n (w_j(t_i) \otimes \tilde{y}_j) \oplus (\tilde{x}_i \ominus_H \sum_{j=1}^n w_j(t_i) \otimes \tilde{x}_j) \otimes a_1 \\ &= \sum_{j=1}^n \left(w_j(t_i) \otimes (y_j; l_{y_j}, r_{y_j})_{LR} \right) \\ &\oplus \left((x_i; l_{x_i}, r_{x_i})_{LR} \ominus_H \sum_{j=1}^n (w_j(t_i) \otimes (x_j; l_{x_j}, r_{x_j})_{LR}) \right) \otimes a_1 \\ &= \left(\sum_{j=1}^n w_j(t_i) y_j; \sum_{j=1}^n w_j(t_i) l_{y_j}, \sum_{j=1}^n w_j(t_i) r_{y_j} \right)_{LR} \\ &\oplus \left(a_1 \otimes (x_i - \sum_{j=1}^n w_j(t_i) x_j); s a_1 \otimes (l_{x_i} - \sum_{j=1}^n w_j(t_i) l_{x_j}) - (1-s) a_1 \right. \\ &\quad \left. \otimes (r_{x_i} - \sum_{j=1}^n w_j(t_i) r_{x_j}) s a_1 \otimes (r_{x_i} - \sum_{j=1}^n w_j(t_i) r_{x_j}) - (1-s) a_1 \otimes (l_{x_i} - \sum_{j=1}^n w_j(t_i) l_{x_j}) \right)_{LR}\end{aligned}$$

که در آن

$$s = \begin{cases} 1 & a_1 \geq 0, \\ 0 & a_1 < 0. \end{cases}$$

۴. روش خوشه‌بندی فازی

فرض کنید φ مجموعه‌ای از n داده و c یک عدد صحیح بزرگتر از یک باشد. گروه‌بندی φ به c قسمت، به صورت مجموعه‌های دو به دو مجزای $\varphi_1, \dots, \varphi_c$ با خصوصیت $\varphi_1 \cup \dots \cup \varphi_c = \varphi$ و برای هر $i, j \in \{1, 2, \dots, c\}$ که $i \neq j$ ، $\varphi_i \cap \varphi_j = \emptyset$ را یک افراز گویند. تابع نشانگر را به صورت زیر در نظر بگیرید

$$\mu_i(X) = \begin{cases} 1 & X \in \varphi_i, \\ 0 & X \notin \varphi_i, \end{cases} \quad i = 1, 2, \dots, c.$$

بر اساس تابع نشانگر فوق، اگر متغیر X در خوشه‌ی i -ام قرار گیرد، مقدار تابع نشانگر آن یک در غیر این صورت صفر است. یکی از نارسایی‌های روش‌های خوشه‌بندی قطعی این است که در

این روش‌ها هر شی می‌تواند متعلق به یک و فقط یک خوشه باشد، بنابراین هر شی به‌گونه‌ای در یکی از خوشه‌ها جای می‌گیرد، هرچند آن شی با آن خوشه همگون نباشد. برای حل این مشکل می‌توان از روش‌های خوشه‌بندی فازی استفاده کرد.

یانگ و کو (۱۹۹۶) یک روش خوشه‌بندی فازی را ارائه کردند. روش آن‌ها، تعمیمی از روش متداول خوشه‌بندی k -میانگین، برای حالتی است که داده‌ها به‌صورت فازی مشاهده شده باشند. فرض کنید یک متغیر بتواند به بیش از یک خوشه با یک درجه عضویتی تعلق داشته باشد. فرض کنید $\mu_i(X)$ یک تابع عضویت باشد، به‌طوری که برای هر $X \in \varphi$ ، $\sum_{i=1}^c \mu_i(X) = 1$. اگر $\varphi = (\tilde{X}_1, \dots, \tilde{X}_n)$ ، مجموعه‌ای از n عدد فازی باشد، هدف در این روش خوشه‌بندی φ در c خوشه، براساس کمینه سازی تابع هدف زیر است

$$F_m(\mu, \tilde{W}) = \sum_{j=1}^n \sum_{i=1}^c \mu_i(\tilde{X}_j) d_f(\tilde{X}_j, \tilde{W}_i),$$

که در آن، $\tilde{W}_i$ مرکز خوشه i -ام، d_f یک فاصله‌ی مناسب بین هر شی از مرکز خوشه و $\mu_i(\tilde{X}_j)$ درجه‌ی عضویت شی j -ام در خوشه i -ام است. مفهوم پایه‌ای در خوشه‌بندی به روش فازی، درجه‌ی عضویت است که نشان می‌دهد یک شی به چه اندازه به خوشه‌ی مورد نظر تعلق دارد. تحلیل رگرسیون با استفاده از روش‌های خوشه‌بندی فازی بر اساس نگرش **یانگ و کو (۱۹۹۷)** و تحت دو روش زیر مورد بررسی قرار می‌گیرد:

(۱) رگرسیون فازی وزنی دو مرحله‌ای: به منظور یافتن معادله رگرسیونی فازی با داده‌های ناهمگون با استفاده از روش وزنی دو مرحله‌ای، ابتدا مشاهدات مربوط به متغیرهای مستقل و متغیر وابسته را به کمک یکی از روش‌های خوشه‌بندی فازی، خوشه‌بندی کرده و مقادیر درجه عضویت را به دست می‌آوریم. سپس این مقادیر را به عنوان وزن در برآورد رگرسیون کمترین مربعات خطا به‌کار می‌بریم.

(۲) رگرسیون فازی تعمیم یافته‌ی یک مرحله‌ای: در این روش خوشه‌بندی فازی و رگرسیون کمترین مربعات در یک روش یک مرحله‌ای با یکدیگر تلفیق می‌شوند، به‌گونه‌ای که در تابع هدف خوشه-بندی فازی به جای مرکز هر خوشه، معادله‌ی خط رگرسیونی مربوط به آن خوشه قرار می‌گیرد و داده‌ها بر مبنای معادله‌ی خط رگرسیونی که بهترین برازش را به داده‌ها داشته باشد، خوشه‌بندی

می‌شوند. لازم به ذکر است که در این روش، داده‌ها به‌طور همزمان خوشه‌بندی شده و معادله‌ی خط رگرسیونی برای هر خوشه به‌دست می‌آید.

فرض کنید مدل رگرسیونی به‌صورت زیر در نظر گرفته شود

$$\tilde{y} = a_0 \oplus a_1 \otimes \tilde{x}.$$

برای برآورد ضرایب مدل رگرسیونی فوق، بر اساس روش تعمیم‌یافته‌ی یک مرحله‌ای و با استفاده از متر d ، تابع هدف به‌صورت زیر تعریف می‌گردد

$$H(a_{0i}, a_{1i}) = \sum_{i=1}^c \sum_{j=1}^n \mu_{ij} d(a_{0i} \oplus a_{1i} \otimes \tilde{x}_j, \tilde{y}_j),$$

که مشابه تابع هدف روش وزنی دو مرحله‌ای است، با این تفاوت که در تابع هدف روش تعمیم‌یافته یک مرحله‌ای مقادیر μ_{ij} ها مجهولند و به ازای $j = 1, 2, \dots, n$ ، $\sum_{i=1}^c \mu_{ij} = 1$. در واقع در روش وزنی دو مرحله‌ای، مقادیر μ_{ij} از قبل از روی یک روش خوشه‌بندی مشخص می‌گردد ولی در این روش، مقادیر μ_{ij} در حین حل مسأله برآورد خواهند شد. حال مسأله‌ی رگرسیون نیمه پارامتری فازی (رابطه‌ی ۱.۳) بر اساس خوشه‌بندی فازی تحت روش رگرسیون تعمیم‌یافته‌ی یک مرحله‌ای، به‌صورت زیر در نظر می‌گیریم

$$\tilde{y}_j = a_{1i} \otimes \tilde{x}_j \oplus \tilde{f}(t_j) \quad t_j = \frac{j}{n}, \quad j = 1, 2, \dots, n, \quad i = 1, 2, \dots, c.$$

می‌خواهیم برآورد معادلات رگرسیونی را بر اساس c خوشه‌ی فازی، مورد بررسی قرار دهیم. لذا مسأله‌ی بهینه‌سازی زیر را در نظر می‌گیریم.

$$\begin{aligned} & \min \sum_{i=1}^c \sum_{j=1}^n \mu_{ij} K(a_{1i} \otimes \tilde{x}_j \oplus \tilde{f}(t_j), \tilde{y}_j), \\ & = \min \sum_{i=1}^c \sum_{j=1}^n \mu_{ij} \int_0^1 |\hat{y}_{j\alpha} - \tilde{y}_{j\alpha}| d\alpha. \\ & s.t \\ & \sum_{i=1}^c \mu_{ij} = 1 \quad j = 1, \dots, n. \end{aligned}$$

که در آن μ_{ij} مقادیر درجات عضویت مربوط به خوشه‌بندی مشاهدات می‌باشد. در این مقاله جهت برآورد پارامترها و مقادیر درجه عضویت، با مینیم کردن مسأله‌ی بهینه سازی فوق از نرم افزار *Minitab* در مثال‌های عددی استفاده شده است.

۵. معیارهای ارزیابی عملکرد مدل

در رگرسیون فازی یک موضوع مهم، ارزیابی عملکرد مدل می‌باشد. در این بخش، تعدادی معیار نیکویی برازش بر اساس **طاهری و کلکین‌نما (۲۰۱۲)** برای اندازه‌گیری برازش یک مدل رگرسیونی معرفی می‌کنیم.

معیار خطا

۱. خطای نسبی برآورد: فرض کنید $\tilde{y}_i$ و $\hat{y}_i$ به ترتیب متغیر وابسته مشاهده شده و برآورد شده از یک مدل رگرسیونی فازی باشند. تفاضل نسبی بین توابع عضویت بین این دو عدد به صورت زیر تعریف می‌شود:

$$E_1(\tilde{y}_i, \hat{y}_i) = \frac{\int_{S_{\tilde{y}_i} \cup S_{\hat{y}_i}} |\tilde{y}_i(x) - \hat{y}_i(x)| dx}{\int_{S_{\tilde{y}_i}} \tilde{y}_i(x) dx}$$

که در آن $S_{\tilde{y}_i}$ و $S_{\hat{y}_i}$ به ترتیب تکیه‌گاه $\tilde{y}_i$ و $\hat{y}_i$ هستند.

۲. خطای برآورد: در تعریف خطای نسبی برآورد اگر فقط تفاضل بین توابع عضویت $\tilde{y}_i$ و $\hat{y}_i$ را در نظر بگیریم و مجموع تابع عضویت عدد فازی مشاهده شده را یک در نظر بگیریم، معیار دیگری به نام خطای برآورد بدست می‌آید که در زیر آن را مشاهده می‌کنید:

$$E_2(\tilde{y}_i, \hat{y}_i) = \int_{S_{\tilde{y}_i} \cup S_{\hat{y}_i}} |\tilde{y}_i(x) - \hat{y}_i(x)| dx$$

مقادیری که $E_1(\tilde{y}_i, \hat{y}_i)$ و $E_2(\tilde{y}_i, \hat{y}_i)$ می‌توانند اختیار کنند در بازه $[0, +\infty)$ است که با یک تغییر متغیر به شکل زیر معیارهای $G_1(\tilde{y}_i, \hat{y}_i)$ و $G_2(\tilde{y}_i, \hat{y}_i)$ بدست می‌آید که محدوده تغییر آن‌ها بازه $[0, 1]$ می‌باشد. هر چه این مقادیر به یک نزدیک‌تر باشند نشان از نزدیکی دو تابع

عضویت است.

$$G_1(\tilde{y}_i, \hat{y}_i) = \frac{1}{1 + E_1(\tilde{y}_i, \hat{y}_i)} \quad G_2(\tilde{y}_i, \hat{y}_i) = \frac{1}{1 + E_2(\tilde{y}_i, \hat{y}_i)}$$

۳. معیار دیگری که بر اساس تفاضل مقادیر بالایی و پایینی α -برش‌های اعداد فازی مشاهده شده و برآورد شده بدست می‌آید، به شکل زیر تعریف می‌شود

$$D(\tilde{y}_i, \hat{y}_i) = \frac{1}{2m} \sum_{k=1}^m (|\hat{y}_i^L[\alpha_k] - y_i^L[\alpha_k]| + |\hat{y}_i^U[\alpha_k] - y_i^U[\alpha_k]|)$$

که در آن بازه‌ی $(0, 1]$ (مقادیر α)، به m قسمت تقسیم می‌شود و به‌طور مثال $\hat{y}_i^L[\alpha_k]$ ، i -امین برآورد کران پایین k -امین α -برش $\tilde{y}_i$ می‌باشد.

در عمل به منظور ارزیابی نیکویی برازش مدل، از روابط زیر استفاده می‌کنیم:

$$G_1 = \frac{\sum_{i=1}^n G_1(\tilde{y}_i, \hat{y}_i)}{n} \quad G_2 = \frac{\sum_{i=1}^n G_2(\tilde{y}_i, \hat{y}_i)}{n}$$

$$ME = \frac{\sum_{i=1}^n E_1(\tilde{y}_i, \hat{y}_i)}{n} \quad MD = \frac{\sum_{i=1}^n D(\tilde{y}_i, \hat{y}_i)}{n}$$

۴. اندازه مشابهت: اگر $\tilde{A}$ و $\tilde{B}$ دو عدد فازی باشند، مقدار مشابهت آن‌ها به شکل زیر تعریف می‌شود:

$$S(\tilde{A}, \tilde{B}) = \frac{Card(\tilde{A} \cap \tilde{B})}{Card(\tilde{A} \cup \tilde{B})}$$

که در آن $Card(\tilde{A})$ به‌صورت زیر تعریف می‌شود.

$$Card(\tilde{A}) = \begin{cases} \sum_x \tilde{A}(x) & \text{اگر } x \text{ گسسته باشد} \\ \int_x \tilde{A}(x) dx & \text{اگر } x \text{ پیوسته باشد} \end{cases}$$

هر چه میانگین اندازه مشابهت اعداد فازی برازش شده و اعداد فازی مشاهده شده پاسخ که به صورت $MSM = \frac{S(\tilde{y}_i, \hat{y}_i)}{n}$ تعریف می‌شود به یک نزدیک باشد، مدل بهتر عمل کرده است.

تعریف ۱.۵. میانگین وزنی معیارهای نیکویی برازش به صورت زیر تعریف می‌شوند.

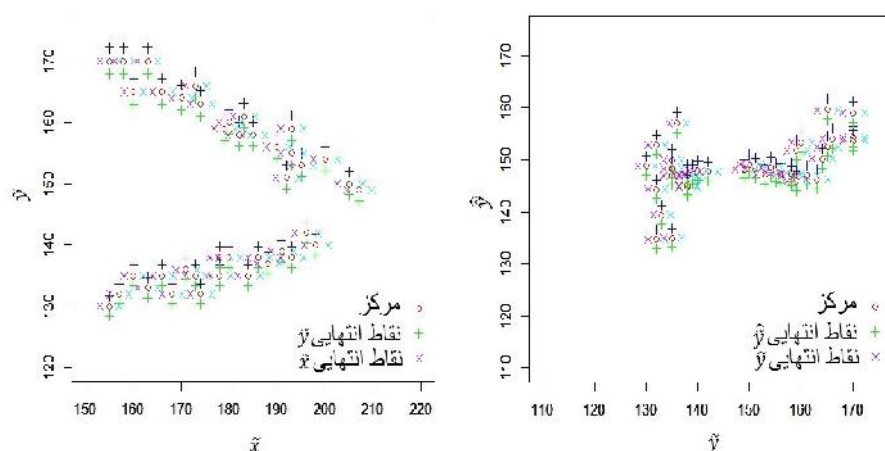
$$G_1(c) = \sum_{i=1}^c w_i G_{1i} \quad G_2(c) = \sum_{i=1}^c w_i G_{2i}$$

$$ME(c) = \sum_{i=1}^c w_i ME_i \quad MD(c) = \sum_{i=1}^c w_i MD_i$$

$$MSM(c) = \sum_{i=1}^c w_i MSM_i$$

که در آن $w_i = \frac{\sum_{j=1}^n \hat{\mu}_{ij}}{\sum_{i=1}^c \sum_{j=1}^n \hat{\mu}_{ij}}$ می‌باشد. در ادامه برای ارزیابی عملکرد مدل‌ها با خوشه‌های متفاوت از این معیارها استفاده می‌کنیم.

مثال ۲.۵. داده‌های فازی جدول ۱ را در نظر بگیرید. مقادیر این جدول شامل وزن (برحسب پوند) و فشار سیستولیک خون (BP) یک نمونه‌ی تصادفی از افراد در گروه‌های سنی ۲۵ تا ۳۰ سال است. لازم به ذکر است که مشاهدات به صورت اعداد فازی مثلثی در نظر گرفته شده‌اند. ابتدا رگرسیون جزئی را با تمام داده‌ها ($c = 1$) انجام داده و معیارهای نیکویی برازش آن را محاسبه می‌کنیم. نتایج در جدول ۲ ارائه شده‌اند. نمودار مربوط به مشاهدات فازی و نمودار $\tilde{y}$



شکل ۱: نمودار مشاهدات و نمودار $\tilde{y}$ در مقابل $\hat{y}$ با $c = 1$ در مثال ۲.۵.

جدول ۱: داده‌های مثال ۲.۵

شماره	وزن	فشار خون	شماره	وزن	فشار خون
	$\tilde{x} = (n; l', r')_T$	$\tilde{y} = (m; l, r)_T$		$\tilde{x} = (n; l', r')_T$	$\tilde{y} = (m; l, r)_T$
۱	(۱۵۵; ۱/۹, ۲)	(۱۳۰; ۱/۶, ۱/۷)	۲۱	(۱۹۸; ۲/۴, ۲/۶)	(۱۴۰; ۱/۷, ۱/۸)
۲	(۱۶۰; ۱/۹, ۲/۱)	(۱۶۵; ۲/۰, ۲/۱)	۲۲	(۲۰۰; ۲/۴, ۲/۶)	(۱۵۴; ۱/۸, ۲/۰)
۳	(۱۶۳; ۲/۰, ۲/۱)	(۱۷۰; ۲/۰, ۲/۲)	۲۳	(۲۰۵; ۲/۵, ۲/۷)	(۱۵۰; ۱/۸, ۲/۰)
۴	(۱۶۸; ۲/۰, ۲/۲)	(۱۳۲; ۱/۶, ۱/۷)	۲۴	(۲۰۷; ۲/۵, ۲/۷)	(۱۴۹; ۱/۸, ۱/۹)
۵	(۱۷۰; ۲/۰, ۲/۲)	(۱۶۴; ۲/۰, ۲/۱)	۲۵	(۱۵۷; ۱/۹, ۲/۰)	(۱۳۲; ۱/۶, ۱/۷)
۶	(۱۷۳; ۲/۱, ۲/۲)	(۱۳۵; ۱/۶, ۱/۸)	۲۶	(۱۶۰; ۱/۹, ۲/۱)	(۱۳۵; ۱/۶, ۱/۸)
۷	(۱۷۴; ۲/۱, ۲/۳)	(۱۶۳; ۲/۰, ۲/۱)	۲۷	(۱۶۳; ۲/۰, ۲/۱)	(۱۳۳; ۱/۶, ۱/۷)
۸	(۱۷۴; ۲/۱, ۲/۳)	(۱۳۲; ۱/۶, ۱/۷)	۲۸	(۱۶۶; ۲/۰, ۲/۲)	(۱۳۵; ۱/۶, ۱/۸)
۹	(۱۷۸; ۲/۱, ۲/۳)	(۱۳۸; ۱/۷, ۱/۸)	۲۹	(۱۵۵; ۱/۹, ۲/۰)	(۱۷۰; ۲/۰, ۲/۲)
۱۰	(۱۷۹; ۲/۱, ۲/۳)	(۱۵۹; ۱/۹, ۲/۱)	۳۰	(۱۵۸; ۱/۹, ۲/۱)	(۱۷۰; ۲/۰, ۲/۲)
۱۱	(۱۸۰; ۲/۲, ۲/۳)	(۱۳۸; ۱/۷, ۱/۸)	۳۱	(۱۶۶; ۲/۰, ۲/۲)	(۱۶۵; ۲/۰, ۲/۱)
۱۲	(۱۸۲; ۲/۲, ۲/۴)	(۱۵۸; ۱/۹, ۲/۱)	۳۲	(۱۷۱; ۲/۱, ۲/۲)	(۱۳۶; ۱/۶, ۱/۸)
۱۳	(۱۸۴; ۲/۲, ۲/۴)	(۱۳۵; ۱/۶, ۱/۸)	۳۳	(۱۷۳; ۲/۱, ۲/۲)	(۱۶۶; ۲/۰, ۲/۲)
۱۴	(۱۸۵; ۲/۲, ۲/۴)	(۱۵۸; ۱/۹, ۲/۱)	۳۴	(۱۸۰; ۲/۲, ۲/۳)	(۱۶۰; ۱/۹, ۲/۱)
۱۵	(۱۸۸; ۲/۳, ۲/۴)	(۱۳۷; ۱/۶, ۱/۸)	۳۵	(۱۷۸; ۲/۱, ۲/۳)	(۱۳۵; ۱/۶, ۱/۸)
۱۶	(۱۹۰; ۲/۳, ۲/۵)	(۱۵۶; ۱/۹, ۲/۰)	۳۶	(۱۸۳; ۲/۲, ۲/۴)	(۱۶۱; ۱/۹, ۲/۱)
۱۷	(۱۹۱; ۲/۳, ۲/۵)	(۱۳۹; ۱/۷, ۱/۸)	۳۷	(۱۸۶; ۲/۲, ۲/۴)	(۱۳۸; ۱/۷, ۱/۸)
۱۸	(۱۹۳; ۲/۳, ۲/۵)	(۱۵۵; ۱/۹, ۲/۰)	۳۸	(۱۹۳; ۲/۳, ۲/۵)	(۱۳۸; ۱/۷, ۱/۸)
۱۹	(۱۹۶; ۲/۴, ۲/۵)	(۱۴۲; ۱/۷, ۱/۸)	۳۹	(۱۹۲; ۲/۳, ۲/۵)	(۱۵۱; ۱/۸, ۲/۰)
۲۰	(۱۹۵; ۲/۳, ۲/۵)	(۱۵۳; ۱/۸, ۲/۰)	۴۰	(۱۹۳; ۲/۳, ۲/۵)	(۱۵۹; ۱/۹, ۲/۱)

در مقابل $\tilde{y}$ در شکل ۱ رسم شده‌اند.

همان‌طور که در شکل ۱ در نمودار $\tilde{y}$ در مقابل $\tilde{x}$ ، واضح است داده‌ها حول نیمساز ربع اول به-گونه‌ای پراکنده‌اند، که می‌توان آن‌ها را به دو دسته تقسیم کرد، که این موضوع، نیاز به خوشه‌بندی داده‌ها به دو دسته را می‌رساند.

حال داده‌ها را به دو دسته خوشه‌بندی کرده و معادله‌ی خط رگرسیونی را برای هر خوشه به‌دست

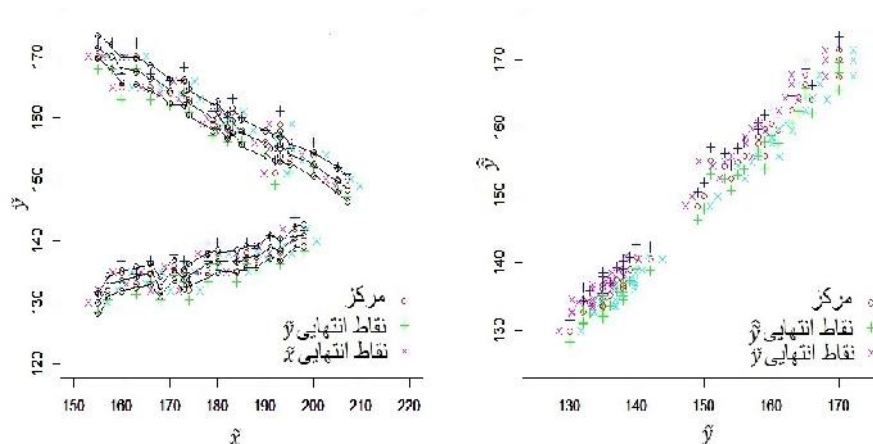
جدول ۲: معیارهای نیکویی برازش با $h = 0.1$ برای حالتیکه $c = 1$

مورد	MSM	G1	G2	ME	MD
مدل پیشنهادی با کرنل سه وزنی $\hat{y} = -0.0223 \otimes \tilde{x} \oplus \tilde{f}(t)$	0.0276	0.3478	0.2260	1.9324	9.3912
مدل یانگ و کو $\hat{y} = 153.2796 \oplus (-0.0271) \otimes \tilde{x}$	0.0004	0.4921	0.3437	1.0321	12.0476

می‌آوریم. نتایج معیارهای نیکویی برازش برای هر خوشه به‌طور جداگانه، به روش رگرسیون نیمه پارامتری با توابع کرنل مختلف و روش یانگ و کو در جدول ۳ آمده است.

در نهایت معیارهای نیکویی برازش برای داده‌های این مثال در جدول ۴ آمده است.

با مقایسه‌ی جدول ۲ با جدول ۴، ملاحظه می‌شود که اگر مدل رگرسیونی را با دو خوشه برآورد



شکل ۲: منحنی برآورد مشاهدات و نمودار $\hat{y}$ در مقابل $\hat{y}$ با $c = 2$ در مثال ۲.۵.

نماییم نتایج بهتری به دست می‌آید و روش پیشنهادی عملکرد بهتری نسبت به روش یانگ و کو داشته است.

جدول ۳: معیارهای نیکویی برازش با $c = ۲$

مورد	h	MSM	$G1$	$G2$	ME	MD
خوشه‌ی اول						
مدل پیشنهادی با کرنل سه وزنی $\tilde{y} = ۰/۱۹۹۶ \otimes \tilde{x} \oplus \tilde{f}(t)$	۰/۱۵۸	۰/۴۰۱۶	۰/۵۵۴۷	۰/۴۳۸۲	۰/۹۳۸۲	۱/۰۰۳۳
مدل پیشنهادی با کرنل مثلثی $\tilde{y} = ۰/۲۲۳۶ \otimes \tilde{x} \oplus \tilde{f}(t)$		۰/۳۸۶۲	۰/۵۴۳۱	۰/۴۲۲۳	۰/۹۶۶۶	۱/۰۴۴۳
مدل پیشنهادی با کرنل نرمال $\tilde{y} = ۰/۲۰۰۵ \otimes \tilde{x} \oplus \tilde{f}(t)$		۰/۳۵۴۲	۰/۵۲۶۶	۰/۴۰۹۱	۱/۰۵۰۴	۱/۱۷۷۸
مدل یانگ و کو $\tilde{y} = ۹۹/۶۱۸۳ \oplus ۰/۲۰۴۷ \otimes \tilde{x}$						
	۰/۱۱۷۴	۰/۵۰۰۹	۰/۳۷۲۱	۱/۰۱۹۹	۱/۴۶۸۱	
خوشه‌ی دوم						
مدل پیشنهادی با کرنل سه وزنی $\tilde{y} = -۰/۴۳۱۳ \otimes \tilde{x} \oplus \tilde{f}(t)$	۰/۱۴۳	۰/۴۲۴۲	۰/۵۷۲۳	۰/۴۲۳۲	۰/۹۲۶۰	۱/۲۹۴۱
مدل پیشنهادی با کرنل مثلثی $\tilde{y} = -۰/۴۶۴۰ \otimes \tilde{x} \oplus \tilde{f}(t)$		۰/۴۴۲۸	۰/۵۸۵۴	۰/۴۴۰۳	۰/۸۹۴۷	۱/۲۳۲۴
مدل پیشنهادی با کرنل نرمال $\tilde{y} = -۰/۴۳۳۹ \otimes \tilde{x} \oplus \tilde{f}(t)$		۰/۴۳۰۳	۰/۵۸۳۳	۰/۴۴۶۷	۰/۹۴۸۵	۱/۴۰۴۹
مدل یانگ و کو $\tilde{y} = ۲۳۲/۸۳۸۶ \oplus (-۰/۴۰۲۹) \otimes \tilde{x}$						
	۰/۲۴۸۹	۰/۵۳۷۰	۰/۳۷۴۳	۰/۹۳۴۲	۱/۶۱۰۶	

جدول ۴: معیارهای نیکویی برازش با $c = ۲$

مدل	$MSM(c = ۲)$	$G1(c = ۲)$	$G2(c = ۲)$	$ME(c = ۲)$	$MD(c = ۲)$
مدل پیشنهادی	۰/۴۲۳۱	۰/۵۷۰۷	۰/۴۳۹۳	۰/۹۱۵۵	۱/۱۲۲۸
مدل یانگ و کو	۰/۱۸۶۰	۰/۵۱۹۷	۰/۳۷۳۲	۰/۹۷۵۲	۱/۵۴۲۴

بحث و نتیجه‌گیری

در این مقاله، چگونگی تعمیم مدل رگرسیون نیمه پارامتری به داده‌های فازی و تحلیل آن با استفاده از خوشه‌بندی فازی بیان شد. همچنین در قالب یک مثال، روش پیشنهادی را با روش یانگ و کو مقایسه کردیم که شاهد عملکرد بهتر روش پیشنهادی بودیم.

مراجع

اسداللهی، م. (۱۳۹۵). برآورد ضرایب مدل‌های رگرسیون خطی بر اساس متغیر وابسته فازی و متغیرهای مستقل فازی و غیر فازی، پایان‌نامه کارشناسی ارشد آمار، دانشکده علوم ریاضی و آمار، دانشگاه بیرجند.

اکبری، م. ق. (۱۳۸۷). استنباط آماری بر اساس داده‌های فازی، رساله دکتری، دانشگاه فردوسی مشهد.

Liu B. (2013), *Uncertainty Theory*, 4th edn, Springer, Berlin.

Robinson P. (1988), Root-N-Consistent Semiparametric Regression, *Econometrica*, 56, 931–954.

Ruspini E. (1969), A new approach to clustering, *Information and Control* 15, 22-32.

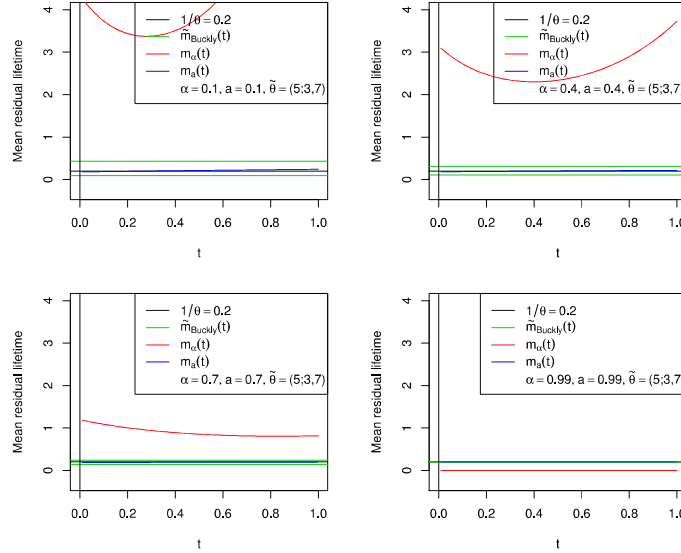
Taheri S. M. and Kelkinnama M. (2012), Fuzzy linear regression based on least absolute deviations, *Iranian Journal of Fuzzy Systems*, 9, 121-140.

Yang M. S. and Ko C. H. (1996), On a class of fuzzy c-numbers clustering procedures for fuzzy data, *Fuzzy Sets and Systems*, 84, 49-60.

Yang M. S. and Ko C. H. (1997), On cluster-wise fuzzy regression analysis, *IEEE Transactions on Systems, Man, and Cybernetics-Part B: Cybernetics*, 27, 1-13.

Zadeh L. A. (1956), Fuzzy sets, *Information and Control* 8, 338-353.

- [3] R. Guo, R. Q. Zhao, D. Guo, and T. Dunne. Random fuzzy variable modeling on repairable system. *Journal of Uncertain Systems*, 1(3): 222–234, 2007.
- [4] R. Guo, R. Q. Zhao, and X. Li. Reliability analysis based on scalar fuzzy variables. *Economic Quality Control*, 22(1):55–70, 2007.
- [5] E. B. Jamkhaneh. An evaluation of the systems reliability using fuzzy lifetime distribution. *Journal of Applied Mathematics, Islamic Azad University of Lahijan*, 7(28):73–80, 2011.
- [6] E. B. Jamkhaneh. Analyzing system reliability using fuzzy weibull lifetime distribution. *International Journal of Applied*, 4(1):93–102, 2014.
- [7] E. B. Jamkhaneh and A. Nozari. Fuzzy system reliability analysis based on confidence interval. In *Advanced Materials Research*, volume 433, pages 4908–4914. Trans Tech Publ, 2012.
- [8] Zdeněk Karpíšek, Petr Štěpánek, and Petr Jurák. Weibull fuzzy probability distribution for reliability of concrete structures. *Engineering mechanics*, 17(5-6):363–372, 2010.
- [9] A. Pak, G. A. Parham, and M. Saraj. Reliability estimation in rayleigh distribution based on fuzzy lifetime data. *International Journal of System Assurance Engineering and Management*, 5(4):487–494, 2014.
- [10] A. R. Saeidi, M. G. Akbari, and M. Doostparast. Hypotheses testing with the two-parameter pareto distribution on the basis of records in fuzzy environment. *Kybernetika*, 50(5):744–757, 2014.
- [11] J. S. Yao, J. S. Su, and T. S. Shih. Fuzzy system reliability analysis using triangular fuzzy numbers based on statistical data. *J. Inf. Sci. Eng.*, 24(5):1521–1535, 2008.

FIGURE 3. Mean residual life function of different approaches based on $\tilde{\theta}$.

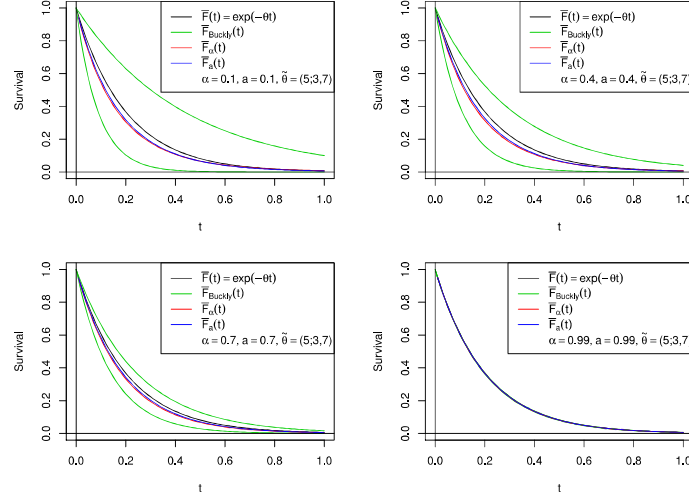
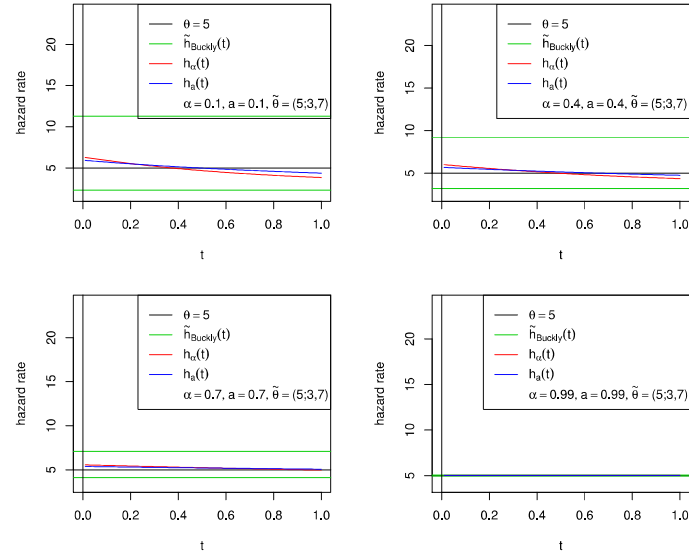
Thus we can say that the best approach is the one based on the density $f_a(t|\tilde{\theta})$.

4. CONCLUSION

In this paper, we provide two new methods based on densities $f_\alpha(t|\tilde{\theta})$ and $f_a(t|\tilde{\theta})$ to consider the performance of reliability function and aging concepts in fuzzy environment. We investigate these methods carefully and compare them with Buckley's approach. An advantage of these new methods is that all functions are single-valued while the Buckley's one is interval-valued. We apply our methods to the situation of exponential life-time random variable and fuzzy number $\tilde{\theta} = (5;3;7)$. After considering the graphs of survival, hazard rate and mean residuals life functions, we observe that the performance of the approach based on density $f_a(t|\tilde{\theta})$ is better than two other methods. Thus we can say that the best methods is the one based on density $f_a(t|\tilde{\theta})$.

REFERENCES

- [1] J. J. Buckley. *Fuzzy probability and statistics*. Springer, 2006.
- [2] H. Garg, M. Rani, and S. P. Sharma. Reliability analysis of the engineering systems using intuitionistic fuzzy set theory. *Journal of Quality and Reliability Engineering*, 2013, 2013.

FIGURE 1. Survival function of different approaches based on $\tilde{\theta}$.FIGURE 2. Hazard rate function of different approaches based on $\tilde{\theta}$.

Consider Figure 2. It shows the hazard rate functions for $\tilde{\theta}$. As we see, in this case all three methods tend to classic mode when α (a for the approach based on $f_a(t|\tilde{\theta})$) tends to 1. It is obvious that the best approach is a new method based on density $f_a(t|\tilde{\theta})$, Because it is the closest to classic mode for each value of a .

Finally Figure 3 shows the performance of the mean residual life functions of different approaches based on $\tilde{\theta}$. It is seen that the approach based on density $f_{\alpha}(t|\tilde{\theta})$ fails to perform well in situation of fuzzy environment.

$$h_{\alpha}(x|\tilde{\theta}) = \frac{\begin{bmatrix} ((\exp(-cx)(c^2x^2 + 2cx + 2))/x^3 \\ -(\exp(-x(c-l+\alpha l))(x^2(c-l+\alpha l)^2 \\ +2x(c-l+\alpha l)+2))/x^3 \\ + (c((\exp(-x(c-l+\alpha l))(x(c-l+\alpha l)+1))/x^2 \\ -(\exp(-cx)(cx+1))/x^2))/l \\ + ((\exp(-cx)(c^2x^2 + 2cx + 2))/x^3 \\ -(\exp(-x(c+u-\alpha u))(x^2(c+u-\alpha u)^2 \\ +2x(c+u-\alpha u)+2))/x^3 \\ - (c((\exp(-cx)(cx+1))/x^2 - (\exp(-x(c+u-\alpha u)) \\ (x(c+u-\alpha u)+1))/x^2))/u \\ + (-\exp(-x(c-l+\alpha l))(x(c-l+\alpha l)+1)) \\ + (\exp(-x(c+u-\alpha u))(x(c+u-\alpha u)+1))/x^2 \end{bmatrix}}{\begin{bmatrix} \left(\exp(-cx-ux) \left(\frac{\exp(\alpha ux) - \exp(ux)}{+ux\exp(ux) - \alpha ux\exp(\alpha ux)} \right) \right) / ux^2 \\ - \left(\exp(-\alpha lx)\exp(-cx) \left(\frac{\exp(\alpha lx) - \exp(lx)}{+lx\exp(\alpha lx) - \alpha lx\exp(lx)} \right) \right) / lx^2 \end{bmatrix}} \quad (2.16)$$

$$m_{\alpha}(x|\tilde{\theta}) = \frac{\int_x^{\infty} \exp(-ct) \left(\begin{array}{l} (\exp((\alpha-1)ut) - 1)/(ut^2) \\ -(1 - \exp((1-\alpha)lt))/(lt^2) \\ -\alpha(\exp((\alpha-1)ut) \\ - \exp(lt(1-\alpha)))/t \end{array} \right) dt}{\exp(-cx) \left(\begin{array}{l} (\exp((\alpha-1)ux) - 1)/(ux^2) \\ -(1 - \exp((1-\alpha)lx))/(lx^2) \\ -\alpha(\exp((\alpha-1)ux) \\ - \exp(lx(1-\alpha)))/x \end{array} \right)} \quad (2.17)$$

3. SIMULATION STUDY

In this section, we suppose that the components of the underlying system have an exponential distribution. Then, we compare the ageing concepts based on various definitions presented in the previous section using the simulation study.

Let X be an exponential random variable whose distribution depends on a fuzzy parameter $\tilde{\theta}$. Our purpose is to compare the ageing concepts based on various approaches for fuzzy density function of exponential distribution. Hence, we consider a triangular fuzzy number $\tilde{\theta} = (5; 3, 7)$ to modeling the vagueness of θ and depict survival, hazard rate and mean residual life functions.

Figure 1 depicts the survival functions of different approaches based on $\tilde{\theta}$. It is seen that the performance of two new approaches is close to each other. The performance of Buckley's approach and approach based on $f_{\alpha}(t|\tilde{\theta})$ are depend on the size of α . It is seen that the former is more sensitive to the value of α than the latter.

$$\bullet \int_{x \in S_X} f_\alpha(x|\tilde{\theta}) dx = 1.$$

Next, we introduce the reliability function and some aging concepts based on this new definition of fuzzy density function.

According to the equation (2.14), one can obtain the reliability, hazard rate and mean residual life functions respectively as follow

$$\begin{aligned} S_\alpha(x|\tilde{\theta}) &= \int_x^\infty f_\alpha(t|\tilde{\theta}) dt \\ &= \frac{\int_x^\infty \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) f(t|\theta) d\theta dt}{\int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) d\theta} \\ &= \frac{\int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) S(x|\theta) d\theta}{\int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) d\theta}, \\ h_\alpha(x|\tilde{\theta}) &= \frac{f_\alpha(x|\tilde{\theta})}{S_\alpha(x|\tilde{\theta})} \\ &= \frac{\int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) f(x|\theta) d\theta}{\int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) S(x|\theta) d\theta}, \\ m_\alpha(x|\tilde{\theta}) &= \frac{\int_x^\infty S_\alpha(t|\tilde{\theta}) dt}{S_\alpha(x|\tilde{\theta})}, \end{aligned}$$

where $f(x|\theta)$ and $S(x|\theta)$ are classical density and survival functions, respectively. For the exponential random variable, the above functions are as follow:

$$S_\alpha(x|\tilde{\theta}) = \frac{- \left[\left(\exp(-cx-ux) \left(\frac{\exp(\alpha ux) - \exp(ux)}{+ux\exp(ux) - \alpha ux\exp(\alpha ux)} \right) \right) / ux^2 \right.}{[(\alpha^2 - 1)(l + u)/2]} \left. - \left(\exp(-\alpha lx)\exp(-cx) \left(\frac{\exp(\alpha lx) - \exp(lx)}{+lx\exp(\alpha lx) - \alpha lx\exp(lx)} \right) \right) / lx^2 \right] \quad (2.15)$$

$$h_a(x|\tilde{\theta}) = \frac{- \left[\begin{aligned} & \left(\begin{aligned} & \exp(-cx)\exp(-ux)(6\exp(ux) - 6\exp(au)) \\ & - 2cx\exp(au) - 2ux\exp(au) + u^2x^2\exp(ux) \\ & + 2cx\exp(ux) - 4ux\exp(ux) - au^2x^2\exp(ux) \\ & + cu^2x^3\exp(ux) + 2aux\exp(ux) + au^2x^2\exp(au) \\ & + 4aux\exp(au) - 2cux^2\exp(ux) - a^2u^2x^2\exp(au) \\ & + acux^2\exp(ux) - acu^2x^3\exp(ux) + acux^2\exp(au) \end{aligned} \right) / (u^2x^4) \\ & + \left(\begin{aligned} & \exp(-alx)\exp(-cx)(6\exp(lx) - 6\exp(alx)) \\ & - l^2x^2\exp(alx) - 2cx\exp(alx) - 4lx\exp(alx) \\ & + 2cx\exp(lx) - 2lx\exp(lx) - al^2x^2\exp(lx) \\ & - 2clx^2\exp(alx) + 4alx\exp(lx) + a^2l^2x^2\exp(lx) \\ & + al^2x^2\exp(alx) - cl^2x^3\exp(alx) + 2alx\exp(alx) \\ & + aclx^2\exp(lx) + aclx^2\exp(alx) + acl^2x^3\exp(alx) \end{aligned} \right) / (l^2x^4) \end{aligned} \right]}{\left[\begin{aligned} & \left(\begin{aligned} & \exp(-alx)\exp(-cx)(2\exp(alx) - 2\exp(lx)) \\ & - l\exp(-alx)\exp(-cx)(ax\exp(alx)) \\ & - 2x\exp(alx) + ax\exp(lx) \end{aligned} \right) / (l^2x^3) \\ & + \left(\begin{aligned} & \exp(-cx - ux)(2\exp(au) - 2\exp(ux)) \\ & - u\exp(-cx - ux)(ax\exp(au)) \\ & - 2x\exp(ux) + ax\exp(ux) \end{aligned} \right) / (u^2x^3) \\ & + \left(\begin{aligned} & \exp(-alx)\exp(-cx)(x^2\exp(alx)) \\ & - ax^2\exp(alx) - (\exp(-cx - ux)(x^2\exp(ux)) \\ & - ax^2\exp(ux)) \end{aligned} \right) / x^3 \end{aligned} \right]} \quad (2.12)$$

$$m_a(x|\tilde{\theta}) = \frac{\int_x^\infty \left(\begin{aligned} & (2\exp(-ct) - 2\exp((1-a)lt - ct)) / (l^2t^3) \\ & + \exp(-ct)(2 - a - a\exp((1-a)lt)) / (lt^2) \\ & + (2\exp((a-1)ut - ct) - 2\exp(-ct)) / (u^2t^3) \\ & + \exp(-ct)(2 - a - a\exp((a-1)ut)) / (ut^2) \end{aligned} \right) dt}{\left[\begin{aligned} & (2\exp(-cx) - 2\exp((1-a)lx - cx)) / (l^2x^3) \\ & + \exp(-cx)(2 - a - a\exp((1-a)lx)) / (lx^2) \\ & + (2\exp((a-1)ux - cx) - 2\exp(-cx)) / (u^2x^3) \\ & + \exp(-cx)(2 - a - a\exp((a-1)ux)) / (ux^2) \end{aligned} \right]} \quad (2.13)$$

2.3. The new density-based approach. Under the general assumptions which mentioned in previous approaches, we define the new density function for fuzzy random variable as

$$f_\alpha(x|\tilde{\theta}) = \frac{\int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) f(x|\theta) d\theta}{\int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) d\theta}, \quad (2.14)$$

where $\mu(\theta)$ is the membership function of $\tilde{\theta}$ and $\tilde{\theta}[\alpha]$ is its α -cut.

It is clear that $f_\alpha(x|\tilde{\theta})$ is the weighted average of $f(x|\theta)$ for the value of $\tilde{\theta}$ in the interval $\tilde{\theta}[\alpha]$ and has the following properties

- $f_\alpha(x|\tilde{\theta}) \geq 0$,

The FPDF of the random variable X is defined by

$$f_a(x|\tilde{\theta}) = \frac{\int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) f(x|\theta) d\theta d\alpha}{\int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) d\theta d\alpha}, \quad a \in [0, 1),$$

where $\mu(\theta)$ is the membership function of the fuzzy parameter $\tilde{\theta}$ and $\tilde{\theta}[\alpha]$ is the corresponding α -cut.

Notice that the above fuzzy density function has the following properties

- $f_a(x|\tilde{\theta}) \geq 0$,
- $\int_{x \in S_X} f_a(x|\tilde{\theta}) dx = 1$.

It is clear that the large value of a makes the proposed FPDF $f_a(x|\tilde{\theta})$ more concentrate on the classical density $f(x|\theta)$.

The following equations, formulate the reliability, hazard rate and mean residual life functions based on this approach, respectively.

$$\begin{aligned} S_a(x|\tilde{\theta}) &= \int_x^\infty f_a(t|\tilde{\theta}) dt \\ &= \frac{\int_x^\infty \int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) f(t|\theta) d\theta d\alpha dt}{\int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) d\theta d\alpha} = \frac{\int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) S(x|\theta) d\theta d\alpha}{\int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) d\theta d\alpha}. \\ h_a(x|\tilde{\theta}) &= \frac{f_a(x|\tilde{\theta})}{S_a(x|\tilde{\theta})} \\ &= \frac{\int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) f(x|\theta) d\theta d\alpha}{\int_x^\infty \int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) f(t|\theta) d\theta d\alpha dt} = \frac{\int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) f(x|\theta) d\theta d\alpha}{\int_a^1 \int_{\theta \in \tilde{\theta}[\alpha]} \mu(\theta) S(x|\theta) d\theta d\alpha}. \\ m_a(x|\tilde{\theta}) &= \frac{\int_x^\infty S_a(t|\tilde{\theta}) dt}{S_a(x|\tilde{\theta})}. \end{aligned}$$

Note that sometimes we have to solve these integrals numerically, since the closed form cannot be achieved.

For the exponential random variable we have:

$$S_a(x|\tilde{\theta}) = \frac{\left[\begin{aligned} &\left(\begin{aligned} &\exp(-alx)\exp(-cx)(2\exp(alx)-2\exp(lx)) \\ &-l\exp(-alx)\exp(-cx)(ax\exp(alx) \\ &-2x\exp(alx)+ax\exp(lx)) \end{aligned} \right) / (l^2x^3) \\ &+ \left(\begin{aligned} &\exp(-cx-ux)(2\exp(aux)-2\exp(ux)) \\ &-u\exp(-cx-ux)(ax\exp(aux) \\ &-2x\exp(ux)+ax\exp(ux)) \end{aligned} \right) / (u^2x^3) \\ &+ \left(\begin{aligned} &\exp(-alx)\exp(-cx)(x^2\exp(alx) \\ &-ax^2\exp(alx)) - (\exp(-cx-ux)(x^2\exp(ux) \\ &-ax^2\exp(ux))) \end{aligned} \right) / x^3 \end{aligned} \right]}{\left[(a-1)^2(a+2)(l+u)/6 \right]} \quad (2.11)$$

$$\begin{aligned}
\tilde{m}(x)[\alpha] &= \left\{ \int_0^\infty S(t|x) dt \mid \theta \in \tilde{\theta}[\alpha] \right\} \\
&= \left\{ \int_0^\infty \frac{S(x+t)}{S(x)} dt \mid \theta \in \tilde{\theta}[\alpha] \right\} \\
&= \left\{ \int_x^\infty \frac{S(u)}{S(x)} du \mid \theta \in \tilde{\theta}[\alpha] \right\}. \tag{2.7}
\end{aligned}$$

Our purpose is to formulate these indexes based on fuzzy exponential lifetime data. Let X be an exponential random variable with the following density function

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

It is easy to show that based on this density the survival, hazard rate and mean residual life functions are respectively obtain as follow

$$S(x; \theta) = e^{-\theta x}, \quad h(x; \theta) = \theta, \quad m(x; \theta) = \frac{1}{\theta}.$$

In order to model the ambiguity of the parameter θ in the underlying problem we use the triangular fuzzy number $\tilde{\theta} = (c; l, u)$. It is obvious that

$$\tilde{\theta}[\alpha] = [\theta_\alpha^l, \theta_\alpha^u] = [c + l(\alpha - 1), c - u(\alpha - 1)], \quad \forall \alpha \in [0, 1].$$

Now, using the relations (2.5-2.7), we will able to obtain the survival, hazard rate and mean residual life functions based on Buckley's approach.

In the case of survival function, since $S(x; \theta)$ is a decreasing function in term of θ , the α -cuts of survival function are given by

$$\tilde{S}(x)[\alpha] = [\exp(-(c - u(\alpha - 1))x), \exp(-(c + l(\alpha - 1))x)] \tag{2.8}$$

Similarly, the α -cuts of the hazard rate and mean residual life functions are obtained respectively as follow,

$$\tilde{h}[\alpha] = [c + l(\alpha - 1), c - u(\alpha - 1)], \tag{2.9}$$

$$\tilde{m}(x)[\alpha] = \left[\frac{1}{c - u(\alpha - 1)}, \frac{1}{c + l(\alpha - 1)} \right]. \tag{2.10}$$

2.2. The Saeidi et al.'s approach. Recently, Saeidi et al. [10] introduced a fuzzy density function in fuzzy environment where the parameter of the distribution was replaced with a fuzzy number. In the following, we recall this approach briefly.

For every $\theta \in \Theta$, we consider a weighted function, say $\mu(\theta)$, for the probability density function $f(x|\theta)$ and define the desired fuzzy probability density function (FPDF), denoted by $f_a(x|\tilde{\theta})$.

where

$$\begin{aligned} f(x)_\alpha^l &= \min \{ f(x|\theta) | \theta \in \tilde{\theta}[\alpha] \}, \\ f(x)_\alpha^u &= \max \{ f(x|\theta) | \theta \in \tilde{\theta}[\alpha] \}. \end{aligned} \quad (2.2)$$

Now, we can obtain the probability of any arbitrary value in the interval $[c, d], c > 0$ based on Equation 2.1. If we denote this probability by $\tilde{P}[c, d]$, then the α -cuts of $\tilde{P}[c, d]$ are defined as (for all $0 \leq \alpha \leq 1$)

$$\tilde{P}[c, d][\alpha] = \left\{ \int_c^d f(x|\theta) dx | \theta \in \tilde{\theta}[\alpha] \right\} = [p_1(\alpha), p_2(\alpha)], \quad (2.3)$$

where

$$\begin{aligned} p_1(\alpha) &= \min \left\{ \int_c^d f(x|\theta) dx | \theta \in \tilde{\theta}[\alpha] \right\}, \\ p_2(\alpha) &= \max \left\{ \int_c^d f(x|\theta) dx | \theta \in \tilde{\theta}[\alpha] \right\}. \end{aligned} \quad (2.4)$$

Then, using relations (2.3-2.4), we can obtain the survival, hazard rate and mean residual life functions as the important ageing concepts. To this end, first we compute the α -cuts of these indexes respectively as (for all $0 \leq \alpha \leq 1$)

$$\begin{aligned} \tilde{S}(x)[\alpha] &= \left[\min \left\{ \int_x^\infty f(t|\theta) dt | \theta \in \tilde{\theta}[\alpha] \right\}, \right. \\ &\quad \left. \max \left\{ \int_x^\infty f(t|\theta) dt | \theta \in \tilde{\theta}[\alpha] \right\} \right], \end{aligned} \quad (2.5)$$

units

$$\begin{aligned} \tilde{h}(x)[\alpha] &= \lim_{\Delta x \rightarrow 0} \frac{\tilde{P}(x < X < x + \Delta x | X > x)[\alpha]}{\Delta x} \\ &= \left\{ \lim_{\Delta x \rightarrow 0} \frac{S(x) - S(x + \Delta x)}{\Delta x S(x)} | \theta \in \tilde{\theta}[\alpha] \right\} \\ &= \left\{ \frac{-S'(x)}{S(x)} | \theta \in \tilde{\theta}[\alpha] \right\} = \left\{ \frac{f(x)}{S(x)} | \theta \in \tilde{\theta}[\alpha] \right\}, \end{aligned} \quad (2.6)$$

[8] defined the fuzzy reliability and its characteristics for the double-state probability model of object based on the fuzzy probability distribution and its properties. Jamkhaneh [5] considered the problem of the evaluation of system reliability, in which the lifetimes of components are described using a fuzzy exponential distribution. Also, Jamkhaneh and Nozari [7] proposed a new method for analyzing the fuzzy system reliability of a parallel-series and series-parallel systems using fuzzy confidence interval, where the reliability of each component of each system is unknown. Garg et al. [2] investigated the reliability analysis of industrial systems by using vague lambda-tau methodology in which information related to system components is uncertain and imprecise in nature. Jamkhaneh [6] proposed a general procedure to construct the reliability characteristics and its α -cut set, when the parameters are fuzzy. Pak et al. [9] presented a Bayesian approach to estimate the parameter and reliability function of Rayleigh distribution from fuzzy lifetime data.

Since the exponential distribution plays an important role in the reliability of a system, we need to extend the classical methods for the uncertainty situation. However, there are many different definitions of the exponential random variable in the uncertainty situation which led to various discussions about the ageing concepts in this situation. In this paper, we propose a new procedure to construct the reliability characteristics of lifetime random variable, when the parameter is fuzzy.

2. AGING CONCEPTS

In this section we want to present the formula (maybe closed or opened form) for some aging concepts such as survival, hazard rate and mean residual life functions based on various approaches in definition of density function of fuzzy exponential distribution.

Thorough this section we suppose that X is a random variable with support $S_X = \{x \in R : f(x|\theta) > 0, \theta \in \Theta\}$, where $f(x; \theta)$ is the density function of X and Θ is the parameter space.

2.1. Buckley's approach. Suppose that due to some conditions in study of underlying problems, we are not able to record the parameter θ precisely. In such situations we use the fuzzy set theory and formulate θ as a fuzzy number $\tilde{\theta}$.

Based on the Buckley's approach, if we assume that $\tilde{\theta}$ is a fuzzy number, then the fuzzy density function is also a fuzzy number with the following α -cuts

$$\tilde{f}(x)[\alpha] = \left[f(x)_\alpha^l, f(x)_\alpha^u \right], \quad (2.1)$$



AGEING CONCEPTS FOR EXPONENTIAL LIFETIME RANDOM VARIABLE IN FUZZY ENVIRONMENT

J. ZENDEHDEL^{1*}, R. ZAREI², M. G. AKBARI¹, AND M. REZAEI¹, .

¹ Department of Statistics, University of Birjand

² Department of Statistics, Faculty of Mathematical Sciences, University of Guilan

ABSTRACT. Ageing concepts are important in reliability engineering. In this paper, we consider the ageing concepts of exponential lifetime random variable in the fuzzy environment. To this end, we consider two density functions which are based on the fuzzy parameter. Using these density functions, we define the survival, the hazard rate and the mean residual life functions.

1. INTRODUCTION

The lifetimes of a component are usually assumed to be random variables with the probability distributions that have crisp parameters. Also, in classical reliability theory, the parameters of lifetime distributions assumed as an exact quantity. However, in the real world, there are situations in which the parameters of the distribution cannot record precisely due to machine errors, experiment, personal judgment or estimation. During the past decades, different approaches and theories have been proposed for treating imprecision and uncertainty, among which the fuzzy set theory is a key one. In the following, we review the main studies on fuzzy reliability in an attempt to display the motivation for this paper.

Buckley [1] introduced fuzzy exponential density using the α -cuts of parameters. Guo et al. [3, 4] investigated random fuzzy variable modeling on repairable system and proposed a credibility hazard concept associated with fuzzy lifetimes based on the conditional distribution of a fuzzy variable under Liu's non-classical credibility measure theory. Yao et al. [11] considered the reliability of serial system and the reliability of parallel system using triangular fuzzy numbers based on statistical data. Karpíšek et al.

Key words and phrases. Reliability engineering, exponential lifetime, Fuzzy parameter, Survival function, Hazard rate function, Mean residual life function.

* speaker.

- Wu J.W., Hung W.L. and Tsai C.H. (2003), Estimation of the parameters of the Gompertz distribution under the first failure-censored sampling plan, *Statistics* 37, 517-525.
- Wang J., Li D. and Chen D. (2012), E-Bayesian Estimation and Hierarchical Bayesian Estimation of the System Reliability Parameter, *Systems Engineering Procedia* 3, 282-289.
- Zadeh L.A. (1986), Probability measures of fuzzy events, *Journal of Mathematical Analysis and Applications* 10, 421-427.

- Al-Hussaini E.K., Al-Dayian G.R. and Adham S.A. (2000), On finite mixture of two-component Gompertz lifetime model, *Journal of Statistical Computation and Simulation* 67, 115.
- Bemmaor A.C. and Glady N. (2012), Modeling purchasing behavior with sudden death: A flexible customer lifetime model, *Management Science* 58(5), 1012-1021.
- Berger J.O. (1985), *Statistical Decision Theory and Bayesian Analysis*, second ed., Springer-Verlag, New York.
- Economos A.C. (1982), Rate of aging, rate of dying and the mechanism of mortality, *Archives of Gerontology and Geriatrics* 1, 46-51.
- Gompertz B. (1825), On the Nature of the Function Expressive of the Law of Human Mortality and on the New Mode of Determining the Value of Life Contingencies, *Philosophical Transactions of the Royal Society American* 115, 513-580.
- Gavrilov L. and Govrilova N. (1991), *The biology of Life Span: A Quantitative approach*, Chur: Harwood.
- Han, M. (1997), The structure of hierarchical prior distribution and its applications, *Chinese Operations Research and Management Science* 63, 31-40.
- Han, M. (2006), E-Bayesian estimation and hierarchical Bayesian estimation of failure rate, *Applied Mathematical Modelling* 33(4), 1915-1922.
- Han, M. (2011), E-Bayesian estimation of the reliability derived from Binomial distribution, *Applied Mathematical Modelling* 35, 2419-2424.
- Iliopoulos G. and Balakrishnan, N. (2011), Exact likelihood inference for Laplace distribution based on Type-II censored samples, *Journal of Statistical Planning and Inference* 141(3), 1224-1239.
- Johnson N.L., Kotz S. and Balakrishnan N. (1995), *Continuous Univariate Distributions*, Vol. 2, 2nd ed., John Wiley and Sons.
- Jaheen, Z. F., Okasha, H. M. (2011), E-Bayesian estimation for the Burr type XII model based on type-2 censoring, *Applied Mathematical Modelling* 35, 4730-4737.
- Kundu D. and Raqab M.Z. (2012), Bayesian inference and prediction of order statistics for a Type-II censored Weibull distribution, *Journal of Statistical Planning and Inference* 142(1), 41-47.
- Lindley D.V. (1980), Approximate Bayesian methods, *Trabajos de Estadística de Investigación Operativa* 31(1), 223-237.
- Pak A., Parham G. H. Saraj, M. (2013), Inference for the Weibull distribution based on fuzzy data, *Revista Colombiana de Estadística* 36(2), 339-358.
- Singh U. and Kumar A. (2007), Bayesian estimation of the exponential parameter under a multiply Type II censoring scheme, *Austrian Journal of Statistics* 36(3), 227-238.
- Smith R. L. and Naylor J. C. (1987), A comparison of maximum likelihood and Bayesian estimators for the three-parameter Weibull distribution, *Applied Statistics* 36, 358-369.

- For fixed n, c , It is observed that as r increases, AV's and MSE's of the E-Bayesian estimation, and the Bayesian estimation is decreases.

3.2. Application with real data set. In this subsection, using a real data set to illustrate the use of the estimation methods proposed in this paper. This data set consists of 63 records of the strengths of 1.5 cm glass fibres by an Institute of Physics in the UK which has been reported in [Smith and Naylor \(1987\)](#). The data are: 0.55, 0.74, 0.77, 0.81, 0.84, 0.93, 1.04, 1.11, 1.13, 1.24, 1.25, 1.27, 1.28, 1.29, 1.30, 1.36, 1.39, 1.42, 1.48, 1.48, 1.49, 1.49, 1.50, 1.50, 1.51, 1.52, 1.53, 1.54, 1.55, 1.55, 1.58, 1.59, 1.60, 1.61, 1.61, 1.61, 1.61, 1.62, 1.62, 1.63, 1.64, 1.66, 1.66, 1.66, 1.67, 1.68, 1.68, 1.69, 1.70, 1.70, 1.73, 1.76, 1.76, 1.77, 1.78, 1.81, 1.82, 1.84, 1.84, 1.89, 2.00, 2.01, 2.24. To compute the Bayesian estimation, E-Bayesian estimation, and hierarchical Bayesian estimation, since we do not have any prior information, we assumed that $a = b = 0.001$. For $a = 3.5$, $c = 2$, $r = 20$, the Bayesian and E-Bayesian estimations become 0.02601762, 0.01601762, respectively. The figures of estimations of the probability density function (1) for the estimation methods are depicted in Figure 1. This represents the superiority of the E-Bayesian estimation toward Bayesian estimation.

4. CONCLUSION

In this study, assuming the scalar parameter of the Gompertz distribution, E-Bayesian estimation of the scalar parameter under the Type II censoring theme was estimated based on fuzzy data with entropy loss function. Then, the Bayesian estimation and E-Bayesian estimation of the scalar parameter of the Gompertz distribution were compared using Monte Carlo simulation. According to the simulation results, the performance of the E-Bayesian estimation is better than that of the Bayesian estimation.

REFERENCES

- Akbari M.G. and Rezaei, A. (2007), An uniformly minimum variance unbiased point estimator using fuzzy observations, *Austrian Journal of Statistics* 36(4), 307-317.

$$X = (X_1, \dots, X_r).$$

$$X = \frac{1}{\beta} \ln \left[1 - \frac{\beta}{\alpha} \ln(1 + U) \right]$$

Each sample of Type II censoring X produced in the third step is considered based on the fuzzy system of [Pak et al. \(2013\)](#) using the membership functions below as a fuzzy sample.

$$\begin{aligned} \mu_{\tilde{x}_1}(x) &= \begin{cases} 1 & x \leq 0/25 \\ \frac{0/5-x}{0/25} & 0/25 \leq x \leq 0/5 \\ 0 & \text{otherwise} \end{cases}, & \mu_{\tilde{x}_2}(x) &= \begin{cases} \frac{x-0/25}{0/25} & 0/25 \leq x \leq 0/5 \\ \frac{0/75-x}{0/25} & 0/5 \leq x \leq 0/75 \\ 0 & \text{otherwise} \end{cases} \\ \mu_{\tilde{x}_3}(x) &= \begin{cases} \frac{x-0/5}{0/25} & 0/5 \leq x \leq 0/75 \\ \frac{1-x}{0/25} & 0/75 \leq x \leq 1 \\ 0 & \text{otherwise} \end{cases}, & \mu_{\tilde{x}_4}(x) &= \begin{cases} \frac{x-0/75}{0/25} & 0/75 \leq x \leq 1 \\ \frac{1/25-x}{0/25} & 1 \leq x \leq 1/25 \\ 0 & \text{otherwise} \end{cases} \\ \mu_{\tilde{x}_5}(x) &= \begin{cases} \frac{x-1}{0/25} & 1 \leq x \leq 1/25 \\ \frac{1/25-x}{0/25} & 1/25 \leq x \leq 1/5 \\ 0 & \text{otherwise} \end{cases}, & \mu_{\tilde{x}_6}(x) &= \begin{cases} \frac{x-1/25}{0/25} & 1/25 \leq x \leq 1/5 \\ \frac{1/75-x}{0/25} & 1/5 \leq x \leq 1/75 \\ 0 & \text{otherwise} \end{cases} \\ \mu_{\tilde{x}_7}(x) &= \begin{cases} \frac{x-1/5}{0/25} & 1/5 \leq x \leq 1/75 \\ \frac{2-x}{0/25} & 1/75 \leq x \leq 2 \\ 0 & \text{otherwise} \end{cases}, & \mu_{\tilde{x}_8}(x) &= \begin{cases} x-1/75 & 1/75 \leq x \leq 2 \\ 1 & x \geq 2 \\ 0 & \text{otherwise} \end{cases} \end{aligned}$$

Then, the Bayesian estimation and E-Bayesian of α were estimated using Equations (9) and (12). Steps one to three have been irritated 5000 times, and the average estimation and its mean square error were estimated and are listed in Tables 1 and 2. Bayesian and E-Bayesian estimates under the entropy loss function is obtained for $\beta = 1.5$ and $c = 2, 3$, respectively. Based on tabulated AV and MSE values, the following conclusions can be drawn from tables.

- The performance of the E-Bayesian estimate of α under entropy loss function is better than that of Bayesian estimation.

TABLE 2. Average value (AV) and mean squared error (MSE) of the estimates of α for different values (n, r) and for $k = 3, \beta = 1.5, c = 3$.

n	r	AV($\hat{\alpha}_{Bay}$)	MSE($\hat{\alpha}_{Bay}$)	AV($\hat{\alpha}_{EBay}$)	MSE($\hat{\alpha}_{EBay}$)
15	10	4.573056530	0.929173595	0.537226004	0.001031632
15	12	2.8398496160	0.0277216463	0.3110839795	0.0004513654
30	10	1.192765392	0.148180635	0.776088361	0.006127616
30	15	1.166947805	0.037062827	0.647970905	0.005298104
30	20	0.6335722291	0.0019095332	0.5005767668	0.0007222743
50	10	2.469167685	0.020121812	0.7616584046	0.001170893
50	20	1.099723789	0.015358997	0.6753051414	0.0007423666
50	30	0.8200900988	0.0040901482	0.634591078	0.000433503
50	40	0.6362921359	0.0001031463	0.618655267	0.0000313089
100	50	0.2742540597	0.0001415297	0.5070088125	0.0004308803
100	60	0.1417788911	0.0000706252	0.4704479375	0.0000380383
100	70	0.0311398726	0.000013259	0.4217555098	0.0000355755

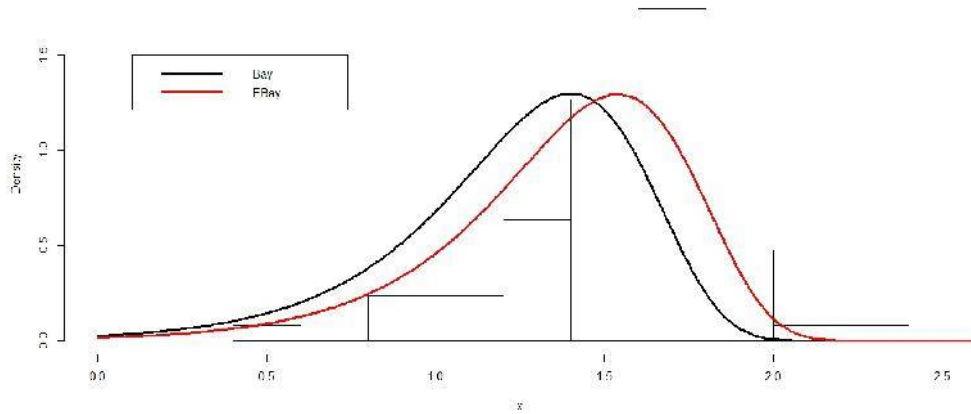


FIGURE 1. Chart probability density function estimation (1) of the E-Bayesian and Bayesian estimations

produced using the equation below where U is a continuous uniform distribution in $(0, 1)$. The produced samples are as

where, δ_{11} and w_3 are given in Equation (10).

3. NUMERICAL EXPERIMENTS

In this section, a numerical example and a Monte Carlo simulation are presented to illustrate all the estimation methods described in the preceding section.

TABLE 1. Average value (AV) and mean squared error (MSE) of the estimates of α for different values (n, r) and for $k = 2, \beta = 1.5, c = 2..$

n	r	AV($\hat{\alpha}_{Bay}$)	MSE($\hat{\alpha}_{Bay}$)	AV($\hat{\alpha}_{EBay}$)	MSE($\hat{\alpha}_{EBay}$)
15	10	2.6389378049	0.0250533562	0.8156070685	0.0009260138
15	12	2.2661401333	0.0206590198	0.6192788289	0.0006851135
30	10	2.2661401333	0.0256590198	0.6192788289	0.0009214633
30	15	0.7508841703	0.0034023036	0.5355348355	0.0006851135
30	20	0.328158153	0.001921408	0.437947491	0.000162547
50	10	1.322396833	0.027452528	0.7329893974	0.002953747
50	20	1.1534889558	0.025642566	0.635105685	0.001747692
50	30	1.120731833	0.014820756	0.591014427	0.001224725
50	40	1.101337432	0.0056261865	0.589955301	0.0009782932
100	50	0.7382615092	0.0010873003	0.6977080349	0.0003964154
100	60	0.6412450135	0.0008641202	0.4456518854	0.0002043431
100	70	0.329805321	0.0000936334	0.406811289	0.0000726658

3.1. Simulation Study. In this subsection, the Bayesian estimation, and E -Bayesian of parameter α are compared with each other. The simulation steps are shown below.

- b is produced per specific c and using the prior distribution of $\pi_1(b) = \frac{1}{c}, 0 < b < c$.
- α is produced using b estimated in the first step and using Equation (5).
- Using the α estimated in the second step and per specific β , the samples of Type II censoring with different (n, r) of Gompertz distribution with a possibility density function (1) are

where,

$$h_1 = \frac{d([h(\alpha)]^{-k})}{d\alpha}, \quad h_{11} = \frac{d^2([h(\alpha)]^{-k})}{d\alpha^2}, \quad v_1 = \frac{dv(\alpha)}{d\alpha},$$

$$w_3 = \frac{d^3w(\alpha)}{d\alpha^3}, \quad \delta_{11} = \left[-\frac{d^2w(\alpha)}{d^2\alpha} \right]^{-1}$$

Given that $h(\alpha) = \alpha$ using Equation (8), the Bayesian estimation of is as follows:

$$\hat{\alpha}_{Bay}(b) = \hat{U}_{Bay}^\alpha(b) = \left(\alpha^{-k} + \frac{k(k+1)\delta_{11}}{2} \alpha^{-(k+2)} + bk\delta_{11} \alpha^{-(k+1)} - \frac{k w_3 (\delta_{11})^2}{2} \right)^{-\frac{1}{k}} \quad (9)$$

where, δ_{11} and w_3 are:

$$\delta_{11} = \left(\frac{n}{\alpha^2} + \frac{1}{\beta^2} \sum_{j=1}^r \frac{I_2^j J_1^j - I_1^j J_2^j}{(J_1^j)^2} \right)^{-1},$$

$$w_3 = \frac{2n}{\alpha^3} + \frac{1}{\beta^3} \sum_{j=1}^r \left(\frac{I_3^j (J_1^j)^2 + I_1 J_1 J_3 - 2I_1^j (J_2^j)^2}{(J_1^j)^3} \right) \quad (10)$$

where,

$$I_i^j = \int_0^\infty (e^{\beta x} - 1) e^{i\beta x - \frac{\alpha}{\beta} e^{\beta x}} \mu_{x_j}(x) dx, \quad J_i^j = \int_0^\infty e^{i\beta x - \frac{\alpha}{\beta} e^{\beta x}} \mu_{x_j}(x) dx, \quad i = 1, 2, 3$$

Definition 2.1. If the Bayesian estimation of parameter θ is $\hat{\theta}_{Bay}(b)$ and the prior distribution of b is $\pi_1(b)$, the E-bayesian estimation of parameter θ that is the mathematical expectation of $\hat{\theta}_{Bay}(b)$ and is shown by $\hat{\theta}_{EBay}$ defined as follows:

$$\hat{\theta}_{EBay} = \int_{\Lambda} \hat{\theta}_{Bay}(b) \pi_1(b) db, \quad b \in \Lambda \quad (11)$$

According to Equations (9) and (11) and the distribution of b , the E-Bayesian estimation of α is defined as follows:

$$\begin{aligned} \hat{\alpha}_{EBay} &= \frac{1}{c} \int_0^c \left(\alpha^{-k} + \frac{k(k+1)\delta_{11}}{2} \alpha^{-(k+2)} + bk\delta_{11} \alpha^{-(k+1)} - \frac{k w_3 (\delta_{11})^2}{2} \right)^{-\frac{1}{k}} db \\ &= \frac{1}{c(k-1)\delta_{11}} \left\{ \left(\alpha^{-k} + \frac{k(k+1)\delta_{11}}{2} \alpha^{-(k+2)} + ck\delta_{11} \alpha^{-(k+1)} - \frac{k w_3 (\delta_{11})^2}{2} \right)^{\frac{k-1}{k}} \right. \\ &\quad \left. - \left(\alpha^{-k} + \frac{k(k+1)\delta_{11}}{2} \alpha^{-(k+2)} - \frac{k w_3 (\delta_{11})^2}{2} \right)^{\frac{k-1}{k}} \right\} \quad (12) \end{aligned}$$

(2011) showed that the most appropriate distribution of b is uniform distribution. Therefore, in this study, the distribution of b was $\pi_1(b)$ as continuous uniform distribution at $(0, c)$, and in this case, Equation (4) is converted as follows:

$$\pi(\alpha|b) = be^{-b\alpha} \quad (5)$$

According to Equations (3) and (5), the prior distribution of providing fuzzy data is as follows:

$$\pi(\alpha|\tilde{\mathbf{x}}) = \frac{\alpha^n e^{-\alpha \left[\frac{(n-r)(e^{\beta m(r)} - 1)}{\beta} \right] + \beta(n-r)m(r)} u(\alpha)}{\int_0^\infty \alpha^n e^{-\alpha \left[\frac{(n-r)(e^{\beta m(r)} - 1)}{\beta} \right] + \beta(n-r)m(r)} u(\alpha) d\alpha} \quad (6)$$

where, $u(\alpha) = \prod_{j=1}^r \left(\int_0^\infty e^{\beta x} e^{-\frac{\alpha}{\beta} e^{\beta x}} \mu_{\tilde{x}_j}(x) dx \right)$. The Bayesian estimation of any function of α , such as $h(\alpha)$, under the entropy loss function is as follows:

$$\begin{aligned} \hat{U}_{Bay}^h(b) &= \{E_{\alpha|\tilde{\mathbf{x}}}(h(\alpha))^{-k}\}^{-\frac{1}{k}} \\ &= \left\{ \frac{\int_0^\infty (h(\alpha))^{-k} \alpha^n e^{-\alpha \left[\frac{(n-r)(e^{\beta m(r)} - 1)}{\beta} \right] + \beta(n-r)m(r)} u(\alpha) d\alpha}{\int_0^\infty \alpha^n e^{-\alpha \left[\frac{(n-r)(e^{\beta m(r)} - 1)}{\beta} \right] + \beta(n-r)m(r)} u(\alpha) d\alpha} \right\}^{-\frac{1}{k}} \end{aligned} \quad (7)$$

In general, it is not possible to calculate it. Therefore, the Lindley (1980) approximation is used to calculate it. Assume:

$$w(\alpha) = \ln l(\tilde{\mathbf{x}}, \alpha) + \ln \pi(b) = L(\alpha) + v(\alpha)$$

where,

$$L(\alpha) = n \ln \alpha - \alpha \left[\frac{(n-r)(e^{\beta m(r)} - 1)}{\beta} \right] + \beta(n-r)m(r) + \ln u(\alpha)$$

and $v(\alpha) = \ln b - b\alpha$. So, Equation (7) is converted as follows.

$$\hat{U}_{Bay}^h(b) = \left\{ \frac{\int_0^\infty [h(\alpha)]^{-k} e^{w(\alpha)} d\alpha}{\int_0^\infty e^{w(\alpha)} d\alpha} \right\}^{-\frac{1}{k}} = \left([h(\alpha)]^{-k} + \frac{1}{2} h_{11} \delta_{11} + v_1 h_1 \delta_{11} + \frac{1}{2} w_3 (\delta_{11})^2 h_1 \right)^{-\frac{1}{k}} \quad (8)$$

with the membership function of $\mu_{\tilde{x}_1}(x_1), \dots, \mu_{\tilde{x}_r}(x_r)$. Let the maximum value of the means of these fuzzy numbers be $m_{(r)}$. The lifetime of $n - r$ surviving units, which are removed from the test after the r^{th} failure, can be encoded as fuzzy numbers $\tilde{x}_{r+1}, \tilde{x}_{r+2}, \dots, \tilde{x}_n$ with the membership functions

$$\mu_{\tilde{x}_j} = \begin{cases} 0 & x \leq m_{(r)} \\ 1 & x > m_{(r)} \end{cases}$$

The fuzzy data $\tilde{\mathbf{x}} = (\tilde{x}_1, \dots, \tilde{x}_n)$ is thus the vector of observed lifetimes. Then, using Zadeh's definition of the probability of a fuzzy event ([Zadeh \(1986\)](#)), the corresponding observed-data likelihood function can be obtained as

$$l(\tilde{\mathbf{x}}, \alpha) = \alpha^n e^{(n-r)\beta m_{(r)}} e^{\frac{n\alpha}{\beta}} e^{\frac{(n-r)\alpha}{\beta}} e^{\beta m_{(r)}} \left[\prod_{j=1}^r \left(\int_0^\infty e^{\beta x} e^{-\frac{\alpha}{\beta} e^{\beta x}} \mu_{\tilde{x}_j}(x) dx \right) \right] \quad (3)$$

where, $u(\alpha) = \prod_{j=1}^r \left(\int_0^\infty e^{\beta x} e^{-\frac{\alpha}{\beta} e^{\beta x}} \mu_{\tilde{x}_j}(x) dx \right)$.

In this paper, E-Bayesian estimation under the entropy loss functions is described in Section 2. A numerical example and a Monte Carlo simulation are presented in Section 3 for illustrative purposes. Section 4 is the conclusions.

2. ESTIMATING THE E-BAYESIAN OF α PARAMETER

Assume the prior distribution of α is the Gamma distribution of density function as follows:

$$\pi(\alpha|a, b) = \frac{b^a}{\Gamma(a)} \alpha^{a-1} e^{-b\alpha}, \alpha > 0, a, b > 0 \quad (4)$$

According to [Han \(1997\)](#), the super parameters of a and b are such that $\pi(\alpha|a, b)$ decreases to α , according to the following equation:

$$\frac{d\pi(\alpha|a, b)}{d\alpha} = \frac{b^a \alpha^{a-2} e^{-b\alpha}}{\Gamma(a)} ((a-1) - b\alpha)$$

where, $b > 0$ and $0 < a \leq 1$. [Berger \(1985\)](#) showed that increasing b resulted in reducing the efficiency of the Bayesian estimator of α . Thus, this super-parameter b should be bounded from top as $0 < b < c$. [Han](#)

Definition 1.1. If X is a reference set, then a fuzzy set of $\tilde{A}$ is shown in using the membership function of $\mu_{\tilde{A}}(x) \rightarrow [0, 1]$ in which per x in X show $\mu_{\tilde{A}}(x)$ membership degree of x in $\tilde{A}$.

Definition 1.2. Fuzzy set of $\tilde{A}$ in reference set of X is normal when and only when $\sup_{x \in X} \mu_{\tilde{A}}(x) = 1$.

Definition 1.3. The fuzzy set of $\tilde{A}$ in the reference set of X is convex, when and if per $x, y \in X$ and $\lambda \in [0, 1]$ the following equation is established:

$$\mu_{\tilde{A}}(\lambda x + (1 - \lambda)y) \geq \{\mu_{\tilde{A}}(x), \mu_{\tilde{A}}(y)\}$$

Definition 1.4. Assuming $L : R^+ \rightarrow [0, 1]$ and $R : R^+ \rightarrow [0, 1]$ are two continuous functions with the following specifications:

$$L(-x) = L(x), \quad R(-x) = R(x)$$

$$L(0) = R(0) = 1$$

$$\lim_{x \rightarrow \infty} L(x) = \lim_{x \rightarrow \infty} R(x) = 1$$

and L and R in $[0, \infty)$ is subtractive, then the fuzzy number of $\tilde{M}$ is a kind of LR when and if the following equation is established:

$$\mu_{\tilde{M}}(x) = \begin{cases} L\left(\frac{m-x}{\alpha}\right) & x \leq m, \alpha > 0 \\ R\left(\frac{m-x}{\beta}\right) & x \geq m, \beta > 0 \end{cases}$$

where, m is the average of the fuzzy number of $\tilde{M}$, α and β are the left and right bounds. Also, the fuzzy number of LR type is shown by $\tilde{M} = (m, \alpha, \beta)$.

Suppose that n independent units are placed on a life test with the corresponding lifetimes $X_1, X_2, \dots, X_n$. It is assumed that these variables are independent and identically distributed with a probability density function (1). Prior to the experiment, a number $r < n$ is determined, and the experiment is terminated after the r^{th} failure. Now, consider the problem where under the Type II censoring scheme, failure times are not precisely observed and only partial information about them is available in the form of fuzzy numbers $\tilde{x}_i = (m, \alpha_i, \beta_i), i = 1, 2, \dots, r$

estimate its parameters

$$f(x; \alpha, \beta) = \alpha e^{\beta x} e^{-\frac{\alpha}{\beta}(e^{\beta x} - 1)}, \quad x > 0, \alpha, \beta > 0 \quad (1)$$

$$F(x; \alpha, \beta) = 1 - e^{-\frac{\alpha}{\beta}(e^{\beta x} - 1)}, \quad x > 0, \alpha, \beta > 0 \quad (2)$$

using E-Bayesian under Type II censoring schemes and based on fuzzy data using entropy loss function as below:

$$L(\hat{\theta}, \theta) \propto \left(\frac{\hat{\theta}}{\theta} \right)^k - k \ln \left(\frac{\hat{\theta}}{\theta} \right) - 1, \quad k \neq 0$$

E- Bayesian method was introduced by [Han \(1997\)](#). Recently, E-Bayesian and hierarchical Bayesian methods have been used by [Han \(2006, 2009\)](#) to estimate exponential estimation parameter and binomial distribution ratio, by [Jaheen and Okasha \(2011\)](#) for Type 12 reliability function distribution parameter estimation based on Type II progressive censoring samples, and [Wang and Chen \(2012\)](#) to estimate Pascal distribution parameters.

Censoring is a common issue in lifetime experiments and reliability studies; the experimenter may not always obtain complete information on failure times for all experimental units. Data obtained from such experiments is called censored data. One of the most common censoring schemes is Type II (failure) censoring, where the life testing experiment is terminated upon the r^{th} (ris pre-fixed) failure. The Type II censoring theme has been the concern of many authors, some of which, like [Singh and Kumar \(2007\)](#) used it for the Bayesian estimation of exponential distribution parameters, [Iliopoulos and Balakrishnan \(2011\)](#) used it in Laplace distribution, and [Kundu and Raqab \(2012\)](#) used it for Bayesian inference in the Weibull distribution.

Some authors have used fuzzy sets in theory estimation. [Akbari and Rezaei \(2007\)](#) proposed a new method for estimating fuzzy spot for uniformly minimum variance. [Pak et al. \(2013\)](#) conducted wide studies on inferential procedures for lifetime distributions based on fuzzy numbers. Here, some definitions are reviewed.



ESTIMATING E-BAYESIAN OF SCALAR PARAMETER OF GOMPERTZ DISTRIBUTION UNDER TYPE II CENSORING SCHEMES BASED ON FUZZY DATA

SHAHRAM YAGHOOBZADEH SHAHRASTANI*

Department of Statistics, Payame Noor University, Tehran, Iran
yagoubzade@gmail.com

ABSTRACT. In this study, the E-Bayesian of the scalar parameter of a Gompertz distribution under Type II censoring schemes were estimated based on fuzzy data under the entropy loss function and using Lindley approximation and the efficiency of the proposed method was compared with the Bayesian estimator using Monte Carlo simulation.

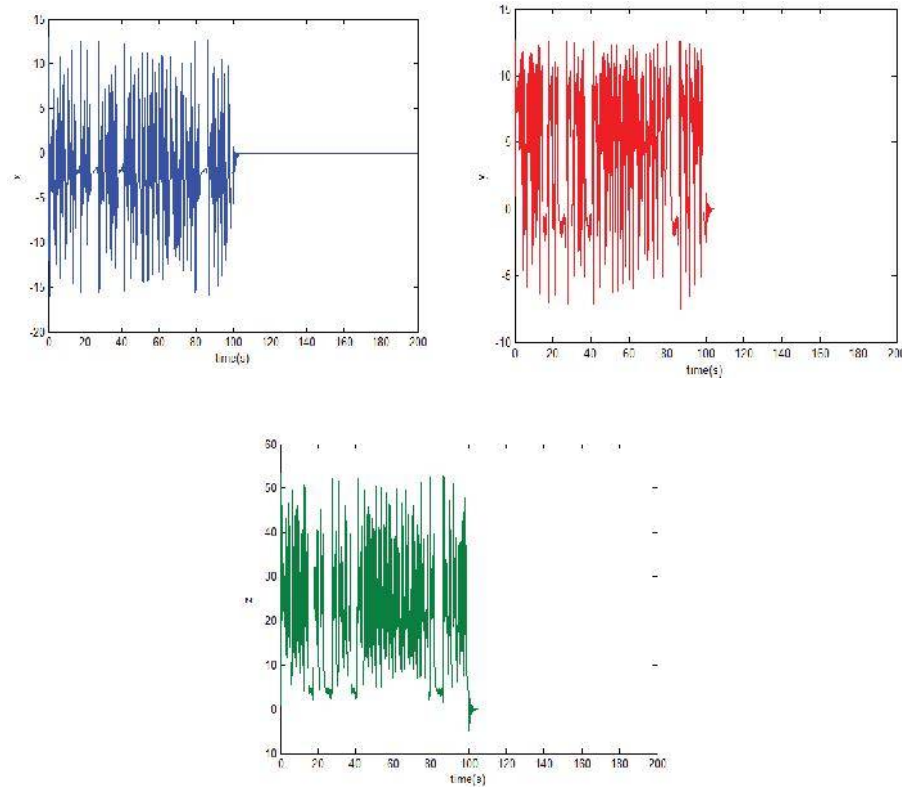
1. INTRODUCTION

The Gompertz distribution, introduced [Gompertz \(1825\)](#), plays an important role in modeling human survival and mortality times ([Gavrilov and Govrilova \(1991\)](#)). It is also applied in various sciences such as biology ([Economos \(1982\)](#)), survival analysis ([Johnson et al. \(1995\)](#)) and marketing ([Bemmaor and Glady \(2012\)](#)). Therefore, estimating its parameters is recommended. Of course, various ways to estimate the distribution parameters of Gompertz has been studied by various authors, including [Al-Hussaini et al. \(2000\)](#), and [Wu et al. \(2003\)](#). In this paper, the probability density function and cumulative distribution function of the Gompertz distribution is considered in order to

Key words and phrases. E-Bayesian Estimation, Gompertz Distribution, Type II Censoring Schemes, Fuzzy Data.

*speaker.

Wang. Kazuo Tanaka, Hua O. (2001), *Fuzzy Control Systems Design and Analysis*, John wiley sons, inc.
Ahmad Taher Azar :2016

FIGURE 4. Stable response of state x , y and z .

REFERENCES

- Ahmad Taher Azar.(2016) , Takagi-Sugeno fuzzy logic controller for Liu-Chen four-scroll chaotic system ,*Int. J. Intelligent Engineering Informatics*, Vol. 4, No. 2.
- Guanrong Chen. David J. Hill. Xinghuo Yu (Eds.). (2003), *Bifurcation Control Theory and Applications*, New York, Springer.
- Liu, W. and Chen, G. (2004) , Can a three-dimensional smooth autonomous quadratic chaotic system generate a single four-scroll chaotic attractor,*International Journal of Bifurcation and Chaos*, Vol. 14, No. 4 , pp. 1395-1403.
- Stephen Wiggins. (2003), *Introduction to Applied Nonlinear Dynamical Systems and Chaos*. Springer-Verlag New York, Inc.
- T. Takagi and M. Sugeno. (1993), Fuzzy Identification of Systems and Its Applications to Modeling and Control, *IEEE Trans. Syst. Man. Cyber.* 15, 116 - 132.

$\lambda_1^2 X^T[(A_1 + B_1 K_1)^T P + P(A_1 + B_1 K_1)]X + \lambda_2^2 X^T[(A_2 + B_2 K_2)^T P + P(A_2 + B_2 K_2)]X + \lambda_1 \lambda_2 X^T[(A_1 + B_1 K_2 + A_2 + B_2 K_1)^T P + P(A_1 + B_1 K_2 + A_2 + B_2 K_1)]X$. Where $V(x) = X^T P X$ and $P = P^T > 0$. When three matrix inequalities $(A_1 + B_1 K_1)^T P + P(A_1 + B_1 K_1) < 0$, $(A_2 + B_2 K_2)^T P + P(A_2 + B_2 K_2) < 0$ and $(A_1 + B_1 K_2 + A_2 + B_2 K_1)^T P + P(A_1 + B_1 K_2 + A_2 + B_2 K_1) < 0$ hold simultaneously, we have $\dot{V}(x) < 0$.

In order to calculate K_i , three linear matrix inequalities $Q A_1^T + A_1 Q + B_1 M_1 + M_1^T B_1^T < 0$, $Q A_2^T + A_2 Q + B_2 M_2 + M_2^T B_2^T < 0$ and $Q(A_1 + A_2)^T + (A_1 + A_2)Q + B_1 M_2 + B_2 M_1 + M_2^T B_1^T + M_1^T B_2^T < 0$ are formed. Where $Q = P^{-1}$, $M_i = K_i Q$, $i=1,2$. The following matrices can be obtained easily. $Q = \begin{bmatrix} 1.2584 & -0.1146 & 0.0831 \\ -0.1146 & 0.2652 & -0.2358 \\ 0.0831 & -0.2358 & 1.3051 \end{bmatrix}$, $M_1 =$

$$\begin{bmatrix} -135.3970 & 27.7985 & 48.3870 \\ -32.2396 & 9.7374 & 24.6706 \\ -51.8636 & 17.3183 & 46.8141 \end{bmatrix}, M_2 = \begin{bmatrix} 49.9288 & 43.4965 & -51.4616 \\ 11.3584 & -32.777 & 77.1689 \\ 56.9448 & -41.5603 & 100.06 \end{bmatrix},$$

$$K_1 = \begin{bmatrix} -101.0472 & 118.9440 & 64.9994 \\ -22.7181 & 53.6075 & 30.0356 \\ -35.7896 & 102.2277 & 58.9184 \end{bmatrix}, K_2 = \begin{bmatrix} 56.6889 & 179.0131 & -10.6903 \\ -1.6409 & -85.3432 & 43.8097 \\ 33.1524 & -90.6468 & 58.1758 \end{bmatrix}.$$

When time $t = 100s$, $u = (\lambda_1 K_1 + \lambda_2 K_2)X$ starts to control the polynomial system and all state variables will converge rapidly to the origin shown as Fig(4).

CONCLUSION

In this paper, we consider Liu-Chen model with a parameter . Bifurcation is discussed simply by employing several numerical simulations on the main parameters. Moreover, the system can be converted into a T-S fuzzy model and a T-S fuzzy control approach is utilized to stabilize the system. Finally, numerical simulations show the effectiveness of the designed T-S fuzzy controller.

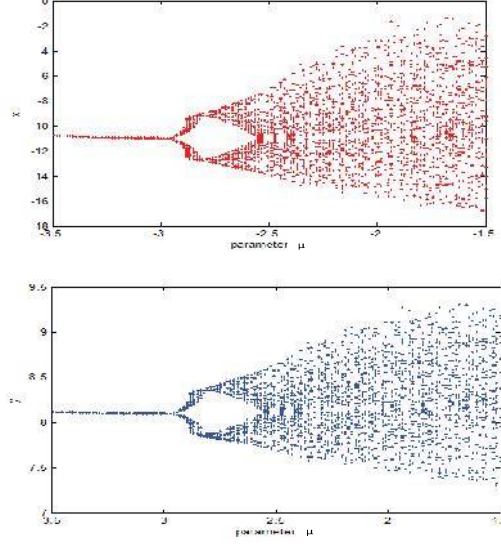


FIGURE 3. Bifurcation of state x and y with the variation of parameter μ .

$$A_1 = \begin{bmatrix} -2 & -5.78 & 7.89 \\ 25.89 & 7.78 & 8 \\ -15.78 & -7.89 & -2 \end{bmatrix}, A_2 = \begin{bmatrix} -2 & 35.48 & -12.74 \\ 5.26 & -33.48 & 8 \\ 25.48 & 12.74 & -2 \end{bmatrix}.$$

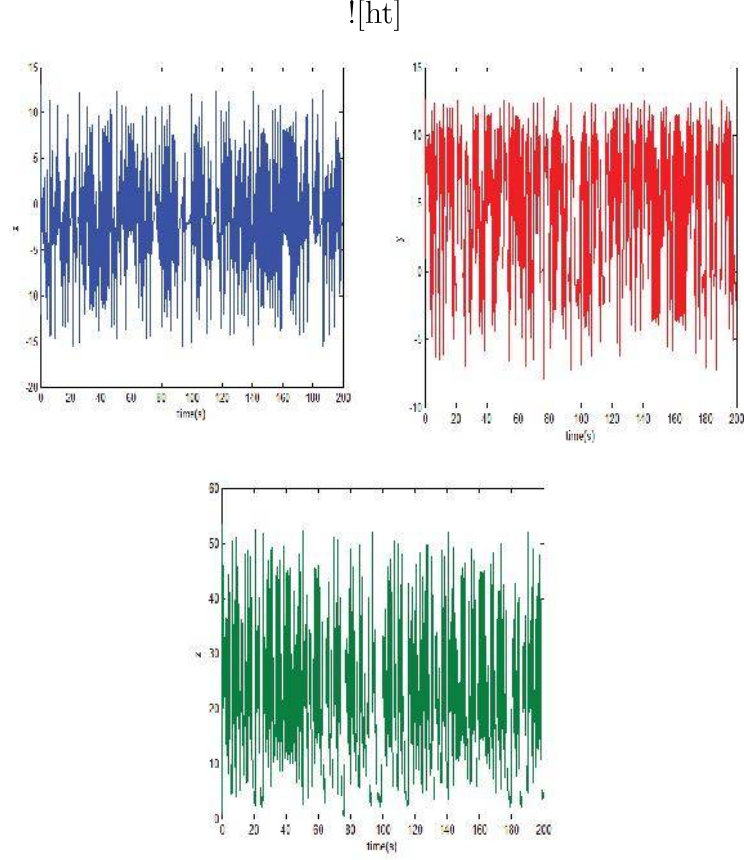
A controlled T-S fuzzy system is described as follows.

IF y is F_1 THEN $\dot{X} = A_1X + B_1u$

IF y is F_2 THEN $\dot{X} = A_2X + B_2u$

$$\text{where } B_1 = \begin{bmatrix} 1 & -1 & -2 \\ 0 & 2 & -1 \\ -1 & 0 & 1 \end{bmatrix}, B_2 = \begin{bmatrix} -1 & 0 & 1 \\ 1 & 1 & 0 \\ -2 & 0 & -1 \end{bmatrix} \text{ and } u = (\lambda_1 K_1 +$$

$\lambda_2 K_2)X$ is a T-S fuzzy controller. The closed-loop system is $\dot{X} = [\lambda_1^2(A_1 + B_1K_1) + \lambda_2^2(A_2 + B_2K_2) + \lambda_1\lambda_2(A_1 + B_1K_2 + A_2 + B_2K_1)]X$ The derivative of a Lyapunov function $V(X)$ to time $\dot{V}(X) = 2X^T P \dot{X} =$

FIGURE 2. Waveform plot of system (1) with $\mu = -2$.

$$\text{where } X = \begin{bmatrix} x \\ y \\ z \end{bmatrix}, A = \begin{bmatrix} -2 & 10 & 0 \\ 18 & -8 & 8 \\ 0 & 0 & -2 \end{bmatrix} \text{ and } B = \begin{bmatrix} 0 & 2 & -1 \\ -1 & -2 & 0 \\ 2 & 1 & 0 \end{bmatrix}.$$

State variable y is bounded. The boundary of y can be estimated simply and precisely within a simulation time interval $t \in [0, 200]$. The range of y is obtained as $y \in [m_1, m_2]$ where $m_1 = -7.89$ and $m_2 = 12.74$. Two membership functions are taken as $\lambda_1 = \frac{m_2 - y}{m_2 - m_1}$, $\lambda_2 = \frac{y - m_1}{m_2 - m_1}$. That $\lambda_1 + \lambda_2 = 1$ and $0 \leq \lambda_i \leq 1$, $i=1,2$.

Matrices A_1 and A_2 are obtained simultaneously.

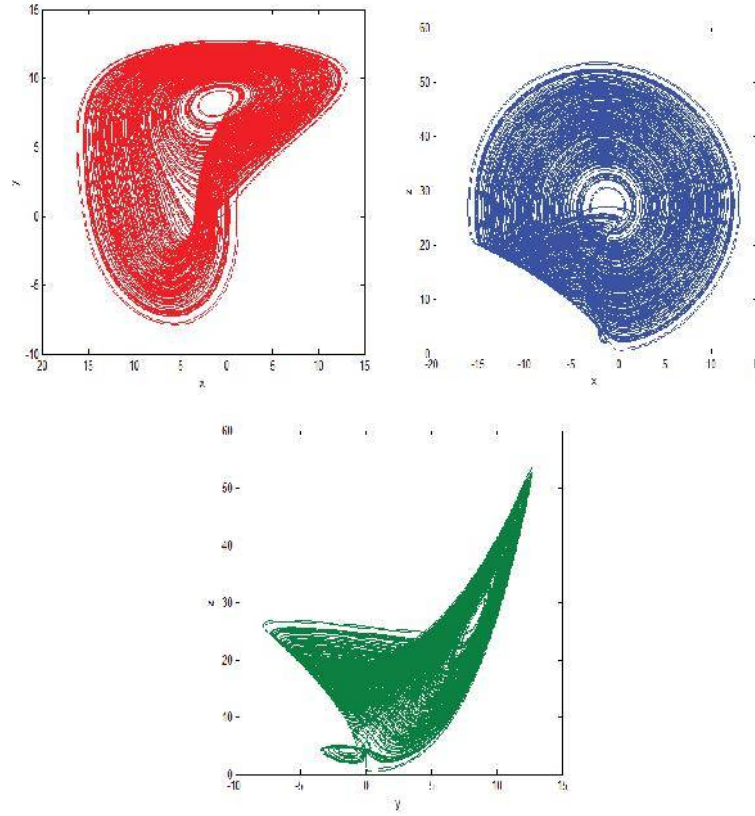


FIGURE 1. Projection of the attractor in x-y, x-z and y-z planes with $\mu = -2$.

rules by analyzing the feature of nonlinearities. These nonlinearities are relevant to a common state variable y . Therefore, a T-S fuzzy model is constructed as follows.

IF y is F_1 THEN $\dot{X} = A_1 X$

IF y is F_2 THEN $\dot{X} = A_2 X$

$$\begin{cases} \dot{x} = \mu x + 10y + 2y^2 - yz \\ \dot{y} = 18x - 8y + 8z - xy - 2y^2 \\ \dot{z} = -2z + 2xy + y^2 \end{cases} \quad (1)$$

The system (1) is described as another form in view of the existence of a common variable in all nonlinear terms

$$\dot{X} = (A + By)X \quad (2)$$

$$\text{where } X = \begin{bmatrix} x \\ y \\ z \end{bmatrix}, A = \begin{bmatrix} \mu & 10 & 0 \\ 18 & -8 & 8 \\ 0 & 0 & -2 \end{bmatrix} \text{ and } B = \begin{bmatrix} 0 & 2 & -1 \\ -1 & -2 & 0 \\ 2 & 1 & 0 \end{bmatrix}.$$

Projections of the attractor in three coordinate planes are shown in Fig(1) and Fig(2).

Fig(1) shows the orbit is around the critical point But it can not be absorbed and Fig(2) shows the critical point is not stable.

Due to the [Stephen Wiggins \(2003\)](#), Bifurcation diagrams of two state variables x and y with the variation of parameter μ are shown respectively in Fig(3). Fig(3) shows that With the passage of $\mu = -2$, changes to the number of critical points. Figures showed that The system has Bifurcation in $\mu = -2$ at origin.

3. T-S FUZZY MODEL AND CONTROL OF THE SYSTEM

In this section, we describe the fuzzy logic controller design of system with using T-S fuzzy model. The stability analysis of the fuzzy logic controller system is carried out using Barbashin-Krasovskii theorem. The system can be converted into a T-S fuzzy model with two simple

of some limit cycles emerging from bifurcation, optimizing the system performance near a bifurcation point, or a combination of some of these objectives. [Chen \(2003\)](#) has collected examples of Bifurcation control.

The [T. Takagi and M. Sugeno \(1985\)](#) proposed a fuzzy model that described by fuzzy IF-THEN rules which represent local linear input-output relations of a nonlinear system. The main feature of a Takagi-Sugeno fuzzy model is to express the local dynamics of each fuzzy implication rule. by a linear system model. The overall fuzzy model of the system is achieved by fuzzy blending of the linear system models. The readers will find that many nonlinear dynamic systems can be represented by Takagi-Sugeno fuzzy models. In fact, it is proved that Takagi-Sugeno fuzzy models are universal approximators. [Wang et al \(2001\)](#) in Their book gives a comprehensive treatment of model-based fuzzy control systems. The central subject of that book is a systematic framework for the stability and design of nonlinear fuzzy control systems. Building on the so-called Takagi-Sugeno fuzzy model, a number of most important issues in fuzzy control systems are addressed. These include stability analysis, systematic design procedures, incorporation of performance specifications, robustness, optimality, numerical implementations, and last but not the least, applications. [Ahmad Taher Azar \(2016\)](#) has presented a fuzzy logic controller design of a chaotic system using TS fuzzy model. The stability analysis of the fuzzy logic controller system has been carried out using Barbashin-Krasovskii theorem.

The organization of this paper is as follows. Analysis of bifurcation in Liu-Chen model are concluded in Section 2. In Section3, T-S fuzzy model of the system is calculated. and a controller for stabilizing the system introduced.

2. EXISTENCE BIFURCATION IN LIU-CHEN MODEL

In this section, we describe the Liu-Chen system ([Liu and Chen \(2004\)](#)) and its qualitative properties. The Liu-Chen system is described by the 3-D dynamics



CONTROL OF BIFURCATION WITH TAKAGI-SUGENO FUZZY LOGIC

MOHAMMAD DARVISHI* AND HAJI MOHAMMAD MOHAMMADINEJAD²

¹*Department of Mathematics, Faculty of Mathematics and Statistics,
University of Birjand
m.darvishi@birjand.ac.ir*

²*Department of Mathematics, Faculty of Mathematics and Statistics,
University of Birjand*

ABSTRACT. paper, a fuzzy logic control using fuzzy sets and fuzzy logic for the control of bifurcation has been proposed. In more detail, this work presents a fuzzy logic controller design of a system using Takagi-Sugeno fuzzy model. The stability analysis of the fuzzy logic controller system is carried out using Barbashin-Krasovskii theorem. For the proposed fuzzy logic control, the Liu-Chen four-scroll system was used. Finally, the simulation results of the proposed control method present the effectiveness of this method.

1. INTRODUCTION

Bifurcation control refers to the task of designing a controller that can modify the bifurcating properties of a given nonlinear system, so as to achieve some desirable dynamical behaviors. Typical bifurcation control objectives include delaying the onset of an inherent bifurcation, introducing a new bifurcation at a preferable parameter value, changing the parameter value of an existing bifurcation point, modifying the shape or type of a bifurcation chain, stabilizing a bifurcated solution or branch, monitoring the multiplicity, amplitude, and/or frequency

Key words and phrases. Takagi-Sugeno fuzzy logic, Control, Bifurcation, Stability.

*speaker.

5. Chen, Y. S. (2001). Outliers detection and confidence interval modification in fuzzy regression, *Fuzzy Sets and Systems*, **119**, 259-272.
6. D'Urso, P. and Massari, R. (2013). Weighted Least Squares and Least Median Squares estimation for the fuzzy linear regression analysis, *Metron* **71** 279-306.
7. D'Urso, P., Massari, R. and Santoro, A. (2011). Robust fuzzy regression analysis, *Information Sciences*, **181**, 4154-4174.
8. Hung, W. L. and Yang, M. S. (2006). An omission approach for detecting outliers in fuzzy regressions models, *Fuzzy Sets and Systems*, **157**, 3109-3122.
9. Leski, J. M. and Kotas, M. (2015). On robust fuzzy c -regression models, *Fuzzy Sets and Systems*, **279**, 112-129.
10. MATLAB, *The Language of Technical Computing*, The MathWorks Inc., MA, 2009.
11. Nguyen, T. D. and Welsch, R. (2010). Outlier detection and least trimmed squares approximation using semi-definite programming, *Comp. Stat. Data Anal.*, **54**, 3212-3226.
12. Rousseeuw, P. J. and Leroy, A. M. (1987). *Robust Regression and Outlier Detection*, Wiley, Hoboken, NJ.
13. Varga, S. (2007). Robust estimations in classical regression models versus robust estimations in fuzzy regression models, *Kybernetika*, **43**, 503-508.
14. Zeng, W., Feng, Q. and Lia, J. (2016). Fuzzy least absolute linear regression, *Applied Soft Computing*, <http://dx.doi.org/10.1016/j.asoc.2016.09.029>.

TABLE 2. Comparison between fuzzy regression models $\widehat{\widehat{y}}_w$ and $\widehat{\widehat{y}}_{LS}$

Models	MSM	MAE_1	MAE_2
$\widehat{\widehat{y}}_w$	0.50	11.09	0.99
$\widehat{\widehat{y}}_{LS}$	0.28	16.02	1.37

the proposed method are 11.09 and 0.99 which are clearly smaller than those of 16.02 and 1.37, respectively, calculated from the LS method. Therefore, like to the conclusion taken by the scatter plot criteria, it can be concluded the estimated weighted fuzzy regression model $\widehat{\widehat{y}}_w$ fits much better than the estimated least-squares fuzzy regression model $\widehat{\widehat{y}}_{LS}$ to the data set.

4. CONCLUSION

The existence of outliers in a set of experimental data can cause incorrect interpretation of the fuzzy linear regression results. The present investigation focused on this problem in a fuzzy regression model for the crisp input and fuzzy output data type and proposed weighted approach to handle the outlier problem. In this regard, a weighted robust fuzzy regression analysis have been developed as an improvement to least-squares estimation in the presence of outliers and to provide us information about what a valid observation is and whether this should be thrown out or down-weighted. The advantages of the proposed approach in contrast with least-squares estimation were compared and discussed in the presence of outliers. This approach not only can detect abnormal values, it can also reduce the impact of these points on the estimation problem by down-weighting them.

REFERENCES

1. Andersen, R. (2007). *Modern Methods for Robust Regression*, Sage, Thousand Oaks, CA.
2. Arefi, M. and Taheri, S. M. (2015). Least-squares regression based on Atanassovs intuitionistic fuzzy inputs-outputs and Atanassovs intuitionistic fuzzy parameters, *IEEE Transactions on Fuzzy Syst.*, **23**, 1142-1154.
3. Chachi, J. and Roozbeh, M. (2017). A fuzzy robust regression approach applied to bedload transport data, *Communications in Statistics-Simulation and Computation*, DOI: 10.1080/03610918.2015.1010002.
4. Chachi, J., Taheri, S. M., Fattahi, S. and Hosseini Ravandi, S. A. (2017). Two robust fuzzy regression models and their application in predicting imperfections of cotton yarn, *Journal of Textiles and Polymers*, In press.

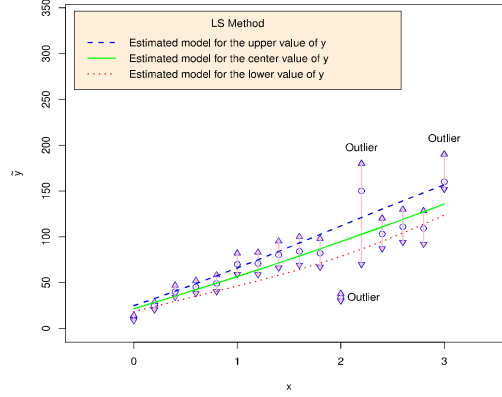


FIGURE 2. The estimated LS model for the data set in Table 1

the presence of a single outlier, irrespective of which element (centers and/or spreads of the response variable, and/or explanatory variable) of the generated data set is contaminated. The robust weighted fit does not suffer from any combination of the given aspects of outliers because such points are down weighted.

3.3. Goodness-of-fit criteria. We employ three criteria known as Mean of Similarity Measures (MSM) and Mean of Absolute Errors (MAE_1 and MAE_2) which have been widely used for evaluation of fuzzy regression models [2, 3, 4]. These criteria are defined as follows:

$$\begin{aligned}
 MSM &= \frac{1}{n} \sum_{i=1}^n \frac{\int \min\{\tilde{y}_i(t), \hat{\tilde{y}}_i(t)\} dt}{\int \max\{\tilde{y}_i(t), \hat{\tilde{y}}_i(t)\} dt}, \\
 MAE_1 &= \frac{1}{n} \sum_{i=1}^n \int |\tilde{y}_i(t) - \hat{\tilde{y}}_i(t)| dt, \\
 MAE_2 &= \frac{1}{n} \sum_{i=1}^n \frac{\int |\tilde{y}_i(t) - \hat{\tilde{y}}_i(t)| dt}{\int \tilde{y}_i(t) dt}.
 \end{aligned}$$

To compare the performances of fuzzy regression models $\hat{\tilde{y}}_w$ and $\hat{\tilde{y}}_{LS}$, the criteria MSM , MAE_1 and MAE_2 are adopted to calculate the accuracy of the models in estimating the observed responses. The corresponding values of these criteria are listed in Table 2 which are almost the same for the two models $\hat{\tilde{y}}_w$ and $\hat{\tilde{y}}_{LS}$. By these results, the MSM value for model $\hat{\tilde{y}}_w$ is 0.50, which is greater than 0.28, the MSM value for model $\hat{\tilde{y}}_{LS}$. On the other hand, the MAE_1 and MAE_2 values for

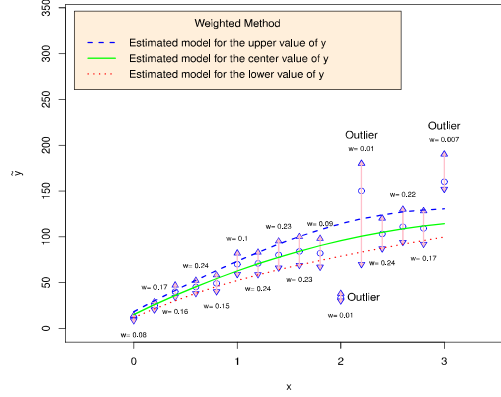


FIGURE 1. The estimated weighted fuzzy regression model for the data set in Table 1

optimal weight values are obtained for these points which are $w_{11} = 0.01$, $w_{12} = 0.01$ and $w_{16} = 0.007$.

Now, in the sequel, in order to provide a comparative study and to investigate the performance of the weighted fuzzy regression model in the presence of outliers a least-squares fuzzy regression model will be considered (called as LS method). In the LS method all of the observations have the same weight $w_i = 1$ for $i = 1, \dots, 16$ in the iteratively reweighted algorithm. Indeed, for the LS method the iteratively reweighted algorithm is an algorithm with fixed values of $w_i = 1$ for $i = 1, \dots, 16$. Now, by applying the LS method to the given data set the following least-squares fuzzy regression model is obtained

$$\begin{aligned} \hat{y}_{LS} = & (21.37 + 33.69x + 1.47x^2, \exp\{1.16 + 1.53x - 0.36x^2\}, \\ & \exp\{1.24 + 1.19x - 0.19x^2\}) \end{aligned}$$

The above model is shown in Fig. 2.

3.2. Scatter plot criteria. D'Urso et al. [7] stated that the methods to detect the presence of outliers relayed on graphical representation and/or analytical procedures. Figures 1 and 2 show the robust weighted and the least-squares fuzzy regression model fits superposed on the scatter plot of the data set considered in Table 1. By analyzing these figures it can be concluded that the least-squares fit is pulled away from the main body of the data by the outlier points. It is clear that the LS-based fuzzy regression model is heavily affected by the presence of outliers in the spreads and/or in the centers of the output. But the weighted fuzzy regression model performance is not affected by

TABLE 1. Data set

i	$(x_i, \tilde{y}_i)$	i	$(x_i, \tilde{y}_i)$
1	$(0.0; (11.5, 3, 2.5)_T)$	9	$(1.6; (84, 15, 16)_T)$
2	$(0.2; (24.8, 4.5, 4)_T)$	10	$(1.8; (82, 15, 16)_T)$
3	$(0.4; (40, 6, 7)_T)$	11	$(2.0; (\mathbf{33}, \mathbf{3}, \mathbf{5})_T)$
4	$(0.6; (45.2, 7, 7)_T)$	12	$(2.2; (\mathbf{150}, \mathbf{80}, \mathbf{30})_T)$
5	$(0.8; (49.1, 9, 9)_T)$	13	$(2.4; (103.1, 16, 17)_T)$
6	$(1.0; (70, 11, 12)_T)$	14	$(2.6; (111, 17, 19)_T)$
7	$(1.2; (70.9, 12, 12)_T)$	15	$(2.8; (109.1, 17, 19)_T)$
8	$(1.4; (80.1, 14, 15)_T)$	16	$(3.0; (\mathbf{160}, \mathbf{8}, \mathbf{30})_T)$

3. OUTLIER DETECTION

This section will exhibit the capability of detecting and down-weighting outliers by the proposed approach and demonstrate how outliers can grossly deteriorate the least-squares estimators of a fuzzy regression model by considering scatter plot criterion and three goodness-of-fit criteria.

3.1. Illustrative example. Consider the data set given in Table 1 which comes from Zeng et al. [14]. The data set consist of crisp input variable x and triangular fuzzy output variable $\tilde{y} = (y, l, r)_T$. Suppose the theoretic model between variables x and $\tilde{y}$ is

$$(y, g(l), h(r)) = \tilde{\beta}_0 \oplus (\tilde{\beta}_1 \otimes x) \oplus (\tilde{\beta}_2 \otimes x^2).$$

We will take $g(\cdot) = h(\cdot) = Ln(\cdot)$ in the solved numerical example. By applying the proposed approach to the given data set in Table 1, the following weighted fuzzy regression model is obtained

$$\begin{aligned} \hat{\tilde{y}}_w = & (15.03 + 54.90x - 7.27x^2, \exp\{1.20 + 1.45x - 0.32x^2\}, \\ & \exp\{1.15 + 1.54x - 0.33x^2\}) \end{aligned}$$

The above model is shown in Fig. 1, where in which the horizontal axis represents the value of the independent variable x and the vertical axis represents the value of the components of triangular fuzzy number, i.e. the upper value $(y + r)$, the center value (y) and the lower value $(y - l)$ of the triangular fuzzy number $\tilde{y} = (y, l, r)_T$. As shown in Fig. 1, the proposed estimation procedure can determine the optimal weight of any observation which are obtained and shown for each point in Fig. 1. From the results given in Fig. 1 it is clear that the $\tilde{y}_{11} = (33, 3, 5)$, $\tilde{y}_{12} = (150, 80, 30)$ and $\tilde{y}_{16} = (160, 8, 30)$ are detected as outliers. According to the estimation procedure the smallest estimated

as follows

$$\begin{aligned}
 (y, g(l), h(r)) &= \tilde{\beta}_0 \oplus (\tilde{\beta}_1 \otimes x_1) \oplus \dots \oplus (\tilde{\beta}_k \otimes x_k) \\
 &= (\beta_0, \sigma_0, \theta_0) \oplus ((\beta_1, \sigma_1, \theta_1) \otimes x) \oplus \dots \\
 &\quad \dots \oplus ((\beta_k, \sigma_k, \theta_k) \otimes x_k) \\
 &= (\beta_0 + \beta_1 x + \dots + \beta_k x_k, \sigma_0 + \sigma_1 |x_1| + \dots + \sigma_k |x_k|, \\
 &\quad \theta_0 + \theta_1 |x_1| + \dots + \theta_k |x_k|).
 \end{aligned}$$

where $g : \mathbb{R}^+ \rightarrow \mathbb{R}$ and $h : \mathbb{R}^+ \rightarrow \mathbb{R}$ are invertible functions. In this model we do not face the non-negativity constraints of the spreads of the response variable because of introducing the invertible functions g and h [6]. Indeed, using this method, the estimated fuzzy response variable is obtained as the following model

$$\begin{aligned}
 \hat{y} &= (\hat{y}, \hat{l}, \hat{r})_{LR} \\
 &= (\hat{\beta}_0 + \hat{\beta}_1 x_1 + \dots + \hat{\beta}_k x_k, g^{-1}\{\hat{\sigma}_0 + \hat{\sigma}_1 |x_1| + \dots + \hat{\sigma}_k |x_k|\}, \\
 &\quad h^{-1}\{\hat{\theta}_0 + \hat{\theta}_1 |x_1| + \dots + \hat{\theta}_k |x_k|\})_{LR}.
 \end{aligned}$$

Theorem 2.1. *In the problem of estimating the fuzzy parameter $\tilde{\beta}$ in following model*

$$(y, g(l), h(r)) = \tilde{\beta}_0 \oplus (\tilde{\beta}_1 \otimes x_1) \oplus \dots \oplus (\tilde{\beta}_k \otimes x_k),$$

the weighted least-squares estimator of $\tilde{\beta}$ is obtained by minimizing the following objective function

$$\begin{aligned}
 E &= \sum_{i=1}^n w_i \mathcal{D}^2 \left((y_i, g(l_i), h(r_i)), \tilde{\beta}_0 \oplus (\tilde{\beta}_1 \otimes x_1) \oplus \dots \oplus (\tilde{\beta}_k \otimes x_k) \right) \\
 &= \sum_{i=1}^n w_i \left[\left(y_i - \sum_{j=0}^k x_{ij} \beta_j \right)^2 + c_1 \left(g(l_i) - \sum_{j=0}^k |x_{ij}| \sigma_j \right)^2 \right. \\
 &\quad \left. + c_2 \left(h(r_i) - \sum_{j=0}^k |x_{ij}| \theta_j \right)^2 \right].
 \end{aligned}$$

In the above objective function $w_i \geq 0$ ($i = 1, \dots, n$) is the weight of the i th residual and must be optimally determined and the constants c_1 and c_2 are defined as $c_1 = \int_0^1 L^{-1}(\alpha) d\alpha$, and $c_2 = \int_0^1 R^{-1}(\alpha) d\alpha$.

Computing parameter estimates and optimal weights in Theorem 2.1 use a method called iteratively reweighted least-squares. which is ran in the software MATLAB [10]. The computing procedure is called iterative because the weights are determined based on the residuals which are changing from iteration to iteration.

are fuzzy (imprecise or vague) [9]. Outliers could be arisen in a fuzzy regression framework with respect to [7]:

- one or more crisp explanatory variables X ;
- the centers and/or the spreads of the fuzzy dependent variable $\tilde{y}$;
- the fuzzy regression model (but not with respect to X and $\tilde{y}$);
- both X and centers of $\tilde{y}$ (but not with respect to the model);
- a combination of the above aspects.

However, in the literature, few approaches are based on using above definition of outliers [3, 6]. Therefore, the main problem which will be investigated in this paper is to provide a suitable robust estimation method to deal with fuzzy data contaminated by the various typologies of outliers described by D'Urso et al. [7]. To this end, a new weighted technique, named as weighted least-squares fuzzy regression, is proposed and applied to establish predictive models in a fuzzy domain. The main idea of the estimation problem is based on detecting and down-weighting unusual data. Because, a data point should not be removed simply because it is an outlier, but it is important to know how these points affect the model fit. Finally, the results of the numerical examples show that weighted approach is able to neutralize and/or smooth the disruptive effects of possible crisp/fuzzy outliers in the estimation process.

The rest of the paper is organized as follows. In the next section the main idea of the paper is explained. Section 3 reports illustrative examples to demonstrate the effectiveness method with the estimates of the regression coefficients of the proposed model in the presence of outliers. In Section 4, some final remarks conclude the paper.

2. A WEIGHTED APPROACH TO FUZZY REGRESSION

Suppose the data set consists of n observations on a response fuzzy-valued variable $\tilde{y}$ and k explanatory real-valued variables

$$\mathbf{x}_i = [x_{i0}, x_{i1}, x_{i2}, \dots, x_{ik}] \in \mathbb{R}^{k+1}, \quad x_{i0} = 1, \quad i = 1, \dots, n,$$

as follows

$$(\tilde{y}_1, \mathbf{x}_1), \dots, (\tilde{y}_n, \mathbf{x}_n),$$

where $\tilde{y}_i = (y_i, l_i, r_i)_{LR}$ ($i = 1, \dots, n$) determines the fuzzy observed of the dependent variable. The relationship between $\tilde{y}$ and $\mathbf{x}$ is formulated



A WEIGHTED LEAST-SQUARES APPROACH FOR THE DETECTION AND DOWN-WEIGHTING OF OUTLIERS FOR FUZZY REGRESSIONS

JALAL CHACHI*

*Department of Mathematics, Statistics and Computer Sciences, Semnan
University
jchachi@semnan.ac.ir*

ABSTRACT. Recent years have seen a surge of interest in extending conventional least-squares fuzzy regression models to robust fuzzy regression model. In this study a weighted fuzzy regression model is proposed dealing with crisp inputs and fuzzy output. The weighted fuzzy regression model is identified by the strategy of optimizing a weighted target function. Some experiments are designed to show its performance in the presence of different kinds of outliers in the data set. The experimental results suggest that the proposed model is capable of dealing with the data set contaminated by outliers and has high prediction accuracy. The proposed fuzzy regression method is capable of determining the weigh of the outlying data points.

1. INTRODUCTION

Robust regression analysis have been developed as an improvement to least-squares estimation in the presence of outliers and to provide us useful information about what a valid observation is and whether this should be thrown out or down-weighted [1, 12, 13]. The primary purpose of robust regression analysis is to fit a model which represents the information in the majority of the data. However, a single outlier affects both on parameters estimates and on the fit of the model to the majority of the data [5, 8, 11]. Therefore, an assessment of potential outliers is important in any analysis especially for formalizing a regression model in a fuzzy domain, in which model parameters and/or data

Key words and phrases. Fuzzy outlier, Fuzzy regression, Weighted least squares.

*speaker.

We use an interactive algorithm to compute α -efficient strategy of the game for Player I so as he/she is satisfied with corresponding security levels.

Generally, in comparison with the existing approaches ([Bigdeli and Hassanpour \(2016\)](#), [Bigdeli et al. \(2016\)](#)), the proposed method has the following advantages:

- (1) This method solves multiobjective matrix game with fuzzy triangular payoffs.
- (2) It does not require any defuzzification.
- (3) Players are satisfied with the obtained solution.
- (4) The proposed approach requires less computational effort.
- (5) It does not require any special order to compare interval inequalities.

4. CONCLUSION

In this paper, a method to find α -efficient strategies and security levels in zero-sum multiobjective game with fuzzy payoffs is presented. It is assumed that fuzzy payoffs are triangular fuzzy numbers which can be presented by their α -cuts. The problem is converted to multiobjective game with interval payoffs. The solutions of this problem are obtained by solving two multiobjective linear programming problems. Then, we use an interactive algorithm for computing α -efficient strategy that player is satisfied with corresponding security level. The main advantage of this method is that it does not require any defuzzification and players are satisfied with the obtained solution.

REFERENCES

- Bector, C. R., Chandra, S. (2005), Fuzzy mathematical programming and Fuzzy matrix games, *Springer Verlag, Berlin*.
- Bigdeli H., Hassanpour H. (2016), A satisfactory strategy of multiobjective two person matrix games with fuzzy payoffs, *Iranian Journal of Fuzzy Systems*, 13, 17-33.
- Bigdeli H., Hassanpour H., Tayyebi J. (2016), The optimistic and pessimistic solutions of single and multiobjective matrix games with fuzzy payoffs and analysis of some of military problems, *Advanced Defence Sci & Tech*, Accepted, (In Persian).
- Nishizaki I., Sakawa M. (2001), Fuzzy and multiobjective games for conflict resolution, *springer-verlag Berlin Heidelberg*.

TABLE 1. The mixed strategies and α -security levels for different values of α in the Example.

α	$x_{1\alpha}^L$	$x_{2\alpha}^L$	$(v_1)_\alpha = [(v_1^L)_\alpha, (v_1^R)_\alpha]$	$(v_2)_\alpha = [(v_2^L)_\alpha, (v_2^R)_\alpha]$
0	0.7916667	0.2083333	[155.2083, 166.3934]	[123.9583, 136.311]
0.1	0.7914573	0.2085427	[155.7927, 165.8317]	[124.5616, 135.6584]
0.2	0.7912458	0.2087542	[156.3771, 165.2557]	[125.1650, 135.0100]
0.3	0.7910321	0.2089679	[156.9615, 164.7258]	[125.7686, 134.3662]
0.4	0.7908163	0.2091837	[157.5459, 164.1818]	[126.3724, 133.7273]
0.5	0.7905983	0.2094017	[158.1303, 163.6441]	[126.9765, 133.0932]
0.6	0.7903780	0.2096220	[158.7148, 163.1126]	[127.5808, 132.4642]
0.7	0.7901554	0.2098446	[159.2992, 162.5876]	[128.1852, 131.8402]
0.8	0.7899306	0.2100694	[159.8837, 162.0692]	[128.7899, 131.2215]
0.9	0.7897033	0.2102967	[160.4682, 161.5575]	[129.3949, 130.6080]
1	0.7894737	0.2105263	[161.0526, 161.0526]	[130, 130]

any α -cut of a triangular fuzzy number is directly obtained from its 1-cut and 0-cut, hence the security levels of Player I are

$$\begin{aligned}\tilde{v}_1 &= ((v_1^L)_0, (v_1^R)_1, (v_1^R)_0) = (155.2083, 161.0526, 166.3934), \\ \tilde{v}_2 &= ((v_2^L)_0, (v_2^R)_1, (v_2^R)_0) = (123.9583, 130, 136.3115).\end{aligned}$$

3.1. Computational result comparsion. In this paper, fuzzy multiobjective matrix game problem led to the two problems with upper and lower bounds of interval objectives. The obtained security levels for Player I are $\tilde{v}_1 = (155.2083, 161.0526, 166.3934)$ and $\tilde{v}_2 = (123.9583, 130, 136.3115)$. To solve presented example in section 3 by the method [Bigdeli and Hassanpour \(2016\)](#), security levels for Player I are $\tilde{v}_1 = (155.2083, 161.0526, 164.6667)$ and $\tilde{v}_2 = (123.9583, 130, 136.0417)$. Note that the solution of the problem in approach of this paper is close to the solution that is reported in method [Bigdeli and Hassanpour \(2016\)](#) (assigned weights by Player I is supposed in both of methodes are equal). The α -efficient strategies of Player I in the method of this paper obtained of solving problem (17) and solutions are close to the obtained α -efficient strategies of method [Bigdeli and Hassanpour \(2016\)](#) for each $\alpha \in [0, 1]$. Observe that, these methods present several strategies for Player I because the fuzzy values are constructed by α -cuts corresponding to different strategies. Decision analyst have to choose one strategy among these strategies to present it to Player I.

where $(\tilde{v}_k)_1 = [(v^{kL})_1, (v^{kR})_1]$ and $(\tilde{v}_k)_0 = [(v^{kL})_0, (v^{kR})_0]$.

Fuzzy value $\tilde{v}_k$ can be expressed as follows:

$$\tilde{v}_k = \bigcup_{\alpha \in [0,1]} (\alpha(\tilde{v}_k)_\alpha) = \bigcup_{\alpha \in [0,1]} (\alpha[\alpha v_k^m + (1-\alpha)v_k^l, \alpha v_k^m + (1-\alpha)v_k^r])$$

which directly implies that the membership function of $\tilde{v}_k$ is

$$\begin{aligned} \mu_{\tilde{v}_k}(x) &= \max \{ \alpha | \alpha v_k^m + (1-\alpha)v_k^l \leq x \leq \alpha v_k^m + (1-\alpha)v_k^r \} \\ &= \begin{cases} \frac{x-v_k^l}{v_k^m-v_k^l} & v_k^l \leq x < v_k^m \\ 1 & x = v_k^m \\ \frac{v_k^r-x}{v_k^r-v_k^m} & v_k^m < x \leq v_k^r \end{cases} \end{aligned}$$

Therefore, the fuzzy value $\tilde{v}_k$ with payoffs of triangular fuzzy numbers is just the triangular fuzzy number $\tilde{v}_k = (v_k^l, v_k^m, v_k^r)$. Also, $\sum_{k=1}^p \lambda_k \tilde{v}_k$ is a triangular fuzzy number.

According to the fact that Player I is cautious, we recommend him to use the obtained strategy by solving the problem (17). This method present several strategies for Player I. Decision analyst have to choose one strategy among these strategies to present it to Player I. In this case, we use the presented interactive algorithm in Bigdeli and Hassanpour (2016) to choose a strategy preferred by Players.

3. NUMERICAL EXAMPLE

Consider the presented example in Bigdeli and Hassanpour (2016) where payoff matrices are

$$\begin{aligned} \tilde{A}_1 &= \begin{bmatrix} (175, 180, 190) & (150, 156, 158) \\ (80, 90, 100) & (175, 180, 190) \end{bmatrix}, \\ \tilde{A}_2 &= \begin{bmatrix} (125, 130, 135) & (120, 130, 135) \\ (120, 130, 140) & (150, 160, 170) \end{bmatrix}. \end{aligned}$$

To find α -efficient strategies and α -cut of security levels of Player I, we solve the two problems (16) and (17) for different value of $\alpha \in [0, 1]$ (suppose $\lambda_1 = 0.5$, $\lambda_2 = 0.5$). The obtained solutions using the simplex method of linear programming are presented in Table 1. Noting that

problems, respectively:

$$\begin{aligned}
& \max \sum_{k=1}^p \lambda_k (\underline{v}^{kR})_\alpha \\
& s.t. \quad \sum_{i=1}^m (a_{ij}^{1R})_\alpha x \geq (\underline{v}^{1R})_\alpha \quad j = 1, \dots, n \\
& \quad \vdots \\
& \quad \sum_{i=1}^m (a_{ij}^{pR})_\alpha x \geq (\underline{v}^{pR})_\alpha \quad j = 1, \dots, n \\
& \quad \sum_{i=1}^m x_i = 1 \\
& \quad x_i \geq 0
\end{aligned} \tag{16}$$

and

$$\begin{aligned}
& \max \sum_{k=1}^k \lambda_k (\underline{v}^{kL})_\alpha \\
& s.t. \quad \sum_{i=1}^m (a_{ij}^{1L})_\alpha x \geq (\underline{v}^{1L})_\alpha \quad j = 1, \dots, n \\
& \quad \vdots \\
& \quad \sum_{i=1}^m (a_{ij}^{pL})_\alpha x \geq (\underline{v}^{pL})_\alpha \quad j = 1, \dots, n \\
& \quad \sum_{i=1}^m x_i = 1 \\
& \quad x_i \geq 0
\end{aligned} \tag{17}$$

Theorem 2.1. Let $\tilde{v}_k, k = 1, \dots, p$ be the fuzzy set defined through the following family of ordinary sets: $\{[(\underline{v}^{kL})_\alpha, (\underline{v}^{kR})_\alpha] \mid \alpha \in [0, 1]\}$ where $(\underline{v}^{kL})_\alpha$ and $(\underline{v}^{kR})_\alpha$ are the obtained solutions of solving the problems (16) and (17) respectively. Then $\tilde{v}_k$ is a fuzzy number.

Proof. Since $\min_{y \in Y} v_\alpha^{kL}(x, y) \leq \min_{y \in Y} v_\alpha^{kR}(x, y)$, for each α , we have $(\underline{v}^{kL})_\alpha \leq (\underline{v}^{kR})_\alpha$. We prove the nesting, i.e, we prove that if $\alpha < \alpha'$, then $[(\underline{v}^{kL})_{\alpha'}, (\underline{v}^{kR})_{\alpha'}] \subseteq [(\underline{v}^{kL})_\alpha, (\underline{v}^{kR})_\alpha]$. It is easy to verify that $(\underline{v}^L)_\alpha \leq (\underline{v}^L)_{\alpha'}$ and $(\underline{v}^R)_{\alpha'} \leq (\underline{v}^R)_\alpha$, for each $\alpha, \alpha' \in [0, 1]$ with $\alpha \leq \alpha'$. This guarantees that the intervals $[(\underline{v}^L)_\alpha, (\underline{v}^R)_\alpha]$ are nested and cosequently, they can be used to construct a fuzzy number. $\square$

Remark 2.2. Note that any α -cut $(\tilde{v}_k)_\alpha$ of the fuzzy value $\tilde{v}_k$ with payoffs of triangular fuzzy numbers can be obtained as

$$(\tilde{v}_k)_\alpha = \alpha(\tilde{v}_k)_1 + (1 - \alpha)(\tilde{v}_k)_0 = [\alpha \underline{v}_k^m + (1 - \alpha) \underline{v}_k^l, \alpha \underline{v}_k^m + (1 - \alpha) \underline{v}_k^r]$$

The problem (13) is equivalent to the following problem which have a finite number of constraints (Bigdeli and Hassanpour (2016)).

$$\begin{aligned}
& \max ((\underline{v}^{1R})_{\alpha}, \dots, (\underline{v}^{pR})_{\alpha}) \\
& s.t. \quad \sum_{i=1}^m (\tilde{a}_{ij}^{1R})_{\alpha} x_i \geq (\underline{v}^{1R})_{\alpha} \quad j = 1, \dots, n \\
& \quad \vdots \\
& \quad \sum_{i=1}^m (\tilde{a}_{ij}^{pR})_{\alpha} x_i \geq (\underline{v}^{pR})_{\alpha} \quad j = 1, \dots, n \\
& \quad \sum_{i=1}^m x_i = 1 \\
& \quad x_i \geq 0 \quad i = 1, \dots, m.
\end{aligned} \tag{14}$$

If Player I regard the problem (11) as pessimistic. It means that Player I supposes that Player II will choose strategy $y \in Y$ such that minimize $v_{\alpha}^{kL}(x, y)$, $k = 1, \dots, p$, then Player I should chooses x so as to maximize $\min v_{\alpha}^{kL}(x, y)$, $k = 1, \dots, p$. Thus we can obtain left extreme points of α -cuts of security levels and corresponding α -efficient strategies to multiobjective game problem by solving the following multiobjective linear programming problem:

$$\begin{aligned}
& \max ((\underline{v}^{1L})_{\alpha}, \dots, (\underline{v}^{pL})_{\alpha}) \\
& s.t. \quad \sum_{i=1}^m (\tilde{a}_{ij}^{1L})_{\alpha} x_i \geq (\underline{v}^{1L})_{\alpha} \quad j = 1, \dots, n \\
& \quad \vdots \\
& \quad \sum_{i=1}^m (\tilde{a}_{ij}^{pL})_{\alpha} x_i \geq (\underline{v}^{pL})_{\alpha} \quad j = 1, \dots, n \\
& \quad \sum_{i=1}^m x_i = 1 \\
& \quad x_i \geq 0 \quad i = 1, \dots, m.
\end{aligned} \tag{15}$$

To solve the problems (14) and (15), we use weighted sum method with $\lambda \in \Lambda^0 = \left\{ \lambda \in \mathbb{R}^p \mid \lambda \geq 0, \sum_{k=1}^p \lambda_k = 1 \right\}$. Thus, we have the following

similar process can be easily introduced for Player II. For a given value of α , Player I has to choose x such that the α -security levels of the game are maximized. Hence, Player I faces with the following multiobjective mathematical programming problem

$$\max_{x \in X} \underline{V}_\alpha(x) \quad (10)$$

or equivalently,

$$\begin{aligned} & \max (\min_{y \in Y} v_\alpha^1(x, y), \dots, \min_{y \in Y} v_\alpha^p(x, y)) \\ \text{s.t. } & \sum_{i=1}^m x_i = 1 \\ & x_i \geq 0 \quad i = 1, \dots, m. \end{aligned} \quad (11)$$

According to the problem (11) for given α , let Player I regard the problem (11) as optimistic. It means that Player I supposes that Player II will choose strategy $y \in Y$ so as to minimize $v_\alpha^{kR}(x, y)$, $k = 1, \dots, p$, then Player I should chooses x so as to maximize $\min v_\alpha^{kR}(x, y)$, $k = 1, \dots, p$, i.e.,

$$\begin{aligned} & \max (\min_{y \in Y} v_\alpha^{1R}(x, y), \dots, \min_{y \in Y} v_\alpha^{pR}(x, y)) \\ \text{s.t. } & \sum_{i=1}^m x_i = 1 \\ & x_i \geq 0 \quad i = 1, \dots, m. \end{aligned} \quad (12)$$

By introducing the auxiliary variables $(\underline{v}^{1R})_\alpha, \dots, (\underline{v}^{pR})_\alpha$, the above problem can be rewritten as

$$\begin{aligned} & \max ((\underline{v}^{1R})_\alpha, \dots, (\underline{v}^{pR})_\alpha) \\ \text{s.t. } & \sum_{j=1}^n \sum_{i=1}^m x_i (\tilde{a}_{ij}^{1R})_\alpha y_j \geq (\underline{v}^{1R})_\alpha \quad \forall y \in Y \\ & \vdots \\ & \sum_{j=1}^n \sum_{i=1}^m x_i (\tilde{a}_{ij}^{pR})_\alpha y_j \geq (\underline{v}^{pR})_\alpha \quad \forall y \in Y \\ & \sum_{i=1}^m x_i = 1 \\ & x_i \geq 0 \quad i = 1, \dots, m. \end{aligned} \quad (13)$$

game with interval payoffs. For each $\alpha \in [0, 1]$, the α -cuts of triangular fuzzy numbers $\tilde{a}_{ij}^k$ are the intervals $(\tilde{a}_{ij}^k)_\alpha = [(\tilde{a}_{ij}^{kL})_\alpha, (\tilde{a}_{ij}^{kR})_\alpha] = [\alpha a_{ij}^{km} + (1 - \alpha)a_{ij}^{kl}, \alpha a_{ij}^{km} + (1 - \alpha)a_{ij}^{kr}]$.

For a fixed value of α , we denote by $\tilde{A}_\alpha^k = [(\tilde{a}_{ij}^k)_\alpha]$, the matrix containing α -cut of the elements of individual matrix $\tilde{A}^k$ for $k = 1, \dots, p$. Assume that we have a fixed value of α . Choosing $x \in X$ and $y \in Y$ by Players I and II, respectively, implies that the expected payoff in level α (α -level expected payoff) of the game is

$$v_\alpha(x, y) = x^T \tilde{A}_\alpha y = [v_\alpha^1(x, y), \dots, v_\alpha^p(x, y)]$$

where $\tilde{A}_\alpha = [\tilde{A}_\alpha^1, \dots, \tilde{A}_\alpha^p]$ and

$$\begin{aligned} v_\alpha^k(x, y) &= x^T \tilde{A}_\alpha^k y = x^T [(\tilde{a}_{ij}^k)_\alpha] y \\ &= [x^T (\tilde{a}_{ij}^{kL})_\alpha y, x^T (\tilde{a}_{ij}^{kR})_\alpha y] \quad k = 1, \dots, p. \end{aligned} \quad (5)$$

Player I (the maximizer) has to find the maximum outcome against any strategy of Player II (the minimizer). Thus, for each strategy $x \in X$ of Player I, the security levels in level α (α -security levels) of Player I are interval payoffs which can be guaranteed against any response of Player II. Therefore, the α -security levels for Players I and II with respect to k -th objective function are defined as follows, respectively:

$$(\underline{V}^k)_\alpha(x) = \min_{y \in Y} v_\alpha^k(x, y) \quad k = 1, \dots, p, \quad (6)$$

$$(\bar{V}^k)_\alpha(y) = \max_{x \in X} v_\alpha^k(x, y) \quad k = 1, \dots, p. \quad (7)$$

For a strategy $x \in X$ and a given value of α , the k -th α -security level of Player I is given by (6), which is a linear programming problem with interval objective function. The α -security levels are p -tuples denoted by

$$\underline{V}_\alpha(x) = ((\underline{V}^1)_\alpha(x), \dots, (\underline{V}^p)_\alpha(x)), \quad (8)$$

$$\bar{V}_\alpha(y) = ((\bar{V}^1)_\alpha(y), \dots, (\bar{V}^p)_\alpha(y)). \quad (9)$$

The above concept allows us to analyze multiobjective matrix games with fuzzy payoffs under the rationale of worst case behavior of the opponent. Here, we state solution process of game for Player I. A

The game defined by (2) is called a two-person zero-sum game with fuzzy payoffs. In this paper, the fuzzy payoffs are considered to be triangular fuzzy numbers $\tilde{a}_{ij} = (a_{ij}^l, a_{ij}^m, a_{ij}^r)$.

For the matrix game $\tilde{A}$ with payoffs of triangular fuzzy numbers, the expected payoff for Player I is computed as

$$\begin{aligned} x^T \tilde{A} y &= \sum_{i=1}^m \sum_{j=1}^n \tilde{a}_{ij} x_i y_j = \\ &= \left(\sum_{i=1}^m \sum_{j=1}^n a_{ij}^l x_i y_j, \sum_{i=1}^m \sum_{j=1}^n a_{ij}^m x_i y_j, \sum_{i=1}^m \sum_{j=1}^n a_{ij}^r x_i y_j \right) \end{aligned}$$

which is a triangular fuzzy number.

Due to the fact that the matrix game $\tilde{A}$ with payoffs of triangular fuzzy numbers is zero-sum, the expected payoff for Player II is obtained as

$$\begin{aligned} x^T (-\tilde{A}) y &= \sum_{i=1}^m \sum_{j=1}^n -\tilde{a}_{ij} x_i y_j = \\ &= \left(\sum_{i=1}^m \sum_{j=1}^n -a_{ij}^r x_i y_j, \sum_{i=1}^m \sum_{j=1}^n -a_{ij}^m x_i y_j, \sum_{i=1}^m \sum_{j=1}^n -a_{ij}^l x_i y_j \right) \end{aligned}$$

which is a triangular fuzzy number too. Thus, in general, Player I's gain floor and Player II's loss ceiling should be triangular fuzzy number. Assume that each player has p objectives. The following multiple fuzzy payoff matrices represent a multiobjective two-person zero-sum game with fuzzy payoffs:

$$\tilde{A}^1 = \begin{bmatrix} \tilde{a}_{11}^1 & \dots & \tilde{a}_{1n}^1 \\ \vdots & \ddots & \vdots \\ \tilde{a}_{m1}^1 & \dots & \tilde{a}_{mn}^1 \end{bmatrix}, \quad \dots, \quad \tilde{A}^p = \begin{bmatrix} \tilde{a}_{11}^p & \dots & \tilde{a}_{1n}^p \\ \vdots & \ddots & \vdots \\ \tilde{a}_{m1}^p & \dots & \tilde{a}_{mn}^p \end{bmatrix}$$

where $\tilde{A}^k$ is the payoff matrix of the game with respect to the k -th objective function, for $k = 1, \dots, p$. The mixed strategy spaces for Players I and II are as follows, respectively

$$X = \{x \in \mathbb{R}^m \mid \sum_{i=1}^m x_i = 1, x_i \geq 0, i = 1, \dots, m\} \quad (3)$$

$$Y = \{y \in \mathbb{R}^n \mid \sum_{j=1}^n y_j = 1, y_j \geq 0, j = 1, \dots, n\} \quad (4)$$

Let each $\tilde{a}_{ij}^k$ be a triangular fuzzy number as $(\tilde{a}_{ij}^{kl}, \tilde{a}_{ij}^{km}, \tilde{a}_{ij}^{kr})$. We use the concept of α -cut to convert the game with fuzzy payoffs to a

for players. First, we recall some definitions of fuzzy sets.

Let X denote a universal set. A fuzzy subset $\tilde{A}$ of X is defined by its membership function $\mu_{\tilde{A}} : X \rightarrow [0, 1]$, which assigns to each element $x \in X$ a real number $\mu_{\tilde{A}}(x)$ in the interval $[0, 1]$. The value of $\mu_{\tilde{A}}(x)$ at x represents the grade of membership of x in $\tilde{A}$. The support of the fuzzy set $\tilde{A}$ on X , denoted by $\text{supp}(\tilde{A})$, is the set of points x in X at which $\mu_{\tilde{A}}(x)$ is positive. The α -cut of the fuzzy set $\tilde{A}$ is defined as ordinary set $\tilde{A}_\alpha = \{x \mid \mu_{\tilde{A}}(x) \geq \alpha\}$, $\alpha \in (0, 1]$, and for $\alpha = 0$, $\tilde{A}_\alpha = \text{closure}\{x \mid \mu_{\tilde{A}}(x) > 0\}$. It plays an important role in the construction of a fuzzy set by a series of ordinary sets. Using the concept of α -cut, the fuzzy set $\tilde{A}$ can be represented by $\tilde{A} = \bigcup_{\alpha \in [0, 1]} \alpha \tilde{A}_\alpha$, where $\alpha \tilde{A}_\alpha$ denotes the algebraic product of a scalar α with the α -cut $\tilde{A}_\alpha$. A triangular fuzzy number $\tilde{A} = (a^l, a^m, a^r)$ is a special fuzzy number, whose membership function is given by

$$\mu_{\tilde{A}}(x) = \begin{cases} (x - a^l)/(a^m - a^l) & a^l \leq x < a^m \\ 1 & x = a^m \\ (a^r - x)/(a^r - a^m) & a^m < x \leq a^r \\ 0 & \text{otherwise,} \end{cases} \quad (1)$$

where a^m is the mean of $\tilde{A}$, and a^l and a^r are the left and right end points of support $\tilde{A}$, respectively.

For any triangular fuzzy number $\tilde{A} = (a^l, a^m, a^r)$, it is easily derived from (1) that $\tilde{A}_\alpha = [a_\alpha^L, a_\alpha^R] = \alpha \tilde{A}_1 + (1 - \alpha) \tilde{A}_0$.

2. TWO-PERSON ZERO-SUM GAME WITH FUZZY PAYOFFS

Consider two players I and II with the set of pure strategies $I = \{1, 2, \dots, m\}$ and $J = \{1, 2, \dots, n\}$, respectively. When Player I chooses a pure strategy $i \in I$ and Player II chooses a pure strategy $j \in J$, let $\tilde{a}_{ij}$ be a fuzzy payoff for Player I and $-\tilde{a}_{ij}$ be a fuzzy payoff for Player II which are assumed to be fuzzy numbers. The two-person zero-sum fuzzy game can be represented by the fuzzy payoff matrix

$$\tilde{A} = \begin{bmatrix} \tilde{a}_{11} & \dots & \tilde{a}_{1n} \\ \vdots & \ddots & \vdots \\ \tilde{a}_{m1} & \dots & \tilde{a}_{mn} \end{bmatrix}. \quad (2)$$



AN APPROACH TO SOLVE MULTIOBJECTIVE MATRIX GAMES WITH FUZZY PAYOFFS

HAMID BIGDELI* AND HASSAN HASSANPOUR

^{1,2}*Department of Mathematics, Faculty of Mathematical Science and
Statistics, University of Birjand
h.bigdeli@birjand.ac.ir*

ABSTRACT. In this paper, an approach is proposed to compute α -efficient strategies and security levels of the players in two-person zero-sum multiobjective game with fuzzy payoffs. It is assumed that fuzzy payoffs are triangular fuzzy numbers. A pair of linear programming models are formulated to obtain the left and right end points of the value of the matrix game in level α . It is shown that the security levels for players are triangular fuzzy numbers. Validity and applicability of the method is illustrated with a practical example.

1. INTRODUCTION

In noncooperative fuzzy games, ambiguity for a player's choice of a strategy, vagueness of preference for a payoff and imprecision of payoff representation are represented by fuzzy sets. Different works are made by several authors in multiobjective matrix games in fuzzy environment (for example see [Bigdeli and Hassanpour \(2016\)](#), [Bigdeli et al. \(2016\)](#), [Nishizaki \(2001\)](#) and in single-objective problems see [Bector and chandra \(2005\)](#)). In this paper, we consider the case of fuzzy matrix games with multiple payoffs in which players' payoffs are expressed as fuzzy numbers and players' pure and mixed strategies are crisp. The presented method computes a satisfactory strategy of game

Key words and phrases. Fuzzy multiobjective game, Interval multiobjective programming, Satisficing α -efficient strategy, Security level.

*speaker.

Contents

An approach to solve multiobjective matrix games with fuzzy payoffs Hamid Bigdeli, Hassan Hassanpour	1
A Weighted Least-Squares Approach for the Detection and ... Jalal Chachi	11
Control of bifurcation with Takagi-Sugeno fuzzy logic Mohammad Darvishi, Haji Mohammad Mohammadinejad	19
Estimating E-Bayesian of Scalar Parameter of Gompertz Distribution ... Shahram Yaghoobzadeh Shahrastani	28
Ageing concepts for exponential lifetime random variable in fuzzy environment J. Zendehdel, R. Zarei, M. G. Akbari, M. Rezaei	40

7th Seminar on Fuzzy Statistics and Probability

University of Birjand

May, 3-4, 2017