

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ



دانشکده علوم ریاضی

رساله دکتری استنباط فازی

برآوردگرهای بهبود یافته در مدل‌های رگرسیون فازی

نگارنده: مجتبی کاشانی

استادان راهنما

دکتر محمد آرشی
دکتر محمدرضا ربیعی

فروردین ۱۴۰۰



دانشگاه شاهرود

مدیریت تحصیلات تکمیلی

شماره:

تاریخ:

باسمه تعالی

ویرایش:

فرم شماره ۶: صورتجلسه دفاع از پایان نامه تحصیلی دوره کارشناسی ارشد

با تأییدات خداوند متعال و با استعانت از حضرت ولی عصر (عج) ارزیابی جلسه دفاع از پایان نامه کارشناسی ارشد
 نام: [نام دانشجو] به شماره دانشجویی [شماره دانشجویی] رشته ریاضی کاربردی گرایش آنالیز عددی تحت عنوان
 [عنوان پایان نامه] همگونی بدون مش کالوگین برای معادلات سهموی تمام غیرمتعلق که در تاریخ ۱۳۹۴/۱۰/۲۷ با حضور هیأت
 محترم داوران در دانشگاه صنعتی شاهرود برگزار گردید به شرح ذیل اعلام می گردد:

قبول (با درجه: بسیار امتیاز ۱۸) ☐ دفاع مجدد ☐ مردود ☐

۲- بسیار خوب (۱۸ - ۱۸/۹۹)

۱- عالی (۲۰ - ۱۹)

۴- قابل قبول (۱۴ - ۱۵/۹۹)

۳- خوب (۱۶ - ۱۷/۹۹)

۵- نمره کمتر از ۱۴ غیر قابل قبول

عضو هیأت داوران	نام و نام خانوادگی	مرتبه علمی	امضاء
۱- استاد راهنمای اول	دکتر علی مس فروش	استادیار	
۲- استاد راهنمای دوم			
۳- استاد مشاور	دکتر مهدی قوتمند	استادیار	
۴- نماینده شورای تحصیلات تکمیلی	دکتر ابراهیم هاشمی	استاد	
۵- استاد ممتحن اول	دکتر حجت احسنی طهرانی	استادیار	
۶- استاد ممتحن دوم	دکتر علیرضا ناظمی	دانشیار	

رئیس دانشکده:

این رساله را با افتخار تقدیم می‌کنم به
همسر عزیزم، بخاطر تمام صبوری‌ها و تشویق‌های
بی‌دریغش که در تمام مسیر زندگی ام گراما بخش
وجودم هست.

همچنین به دختر زیایم مریم و پسر عزیزم
مهدی رضا که وجودشان همچون ستارگانی
در حشان در آسمان زندگی من است.

آن دوست که من دارم آن یار که من دانم
 شیرین دهنی دارد دور از لب و دندانم
 بخت این نکند با من کان شاخ صنوبر را
 بشنیم و بشانم، گل؛ بر سرش افشانم
 عجب سروی، عجب ماهی، عجب باقوت و مرجانی
 عجب جسمی، عجب عقلی، عجب عشقی، عجب جانی
 عجب لطف بهاری تو، عجب میر شکاری تو
 در آن غمزه چه داری تو؟ به زیر لب چه میخوانی؟
 عجب حلوائی قندی تو، امیر بی گزندگی تو
 عجب ماه بلندی تو، که گردون را بگردانی
 تویی کامل، منم ناقص، تویی خالص منم مخلص
 تویی سور و منم راقص، من اسفل تو معلای

سپاس‌گزاری

با عنایت به لطف ایزد یکتا، جا دارد که از اساتید عزیزم آقایان دکتر آرشی و دکتر ربیعی بخاطر تلاش‌های فراوان ایشان در به ثمر نشستن این پروژه و راهنمایی بنده تشکر و قدردانی نمایم. ضمناً تشکر و سپاس فراوان از پرسنل با اخلاق و فهیم دانشکده علوم ریاضی دانشگاه صنعتی شاهرود دارم به واسطه کمک‌ها و همکاری‌های خالصانه ایشان در ایجاد محیطی پر از اطمینان و آرامش برای بنده و بطور قطع برای تمام دانشجویان این دانشکده. از دوستان عزیز دوران تحصیل نیز تشکر ویژه دارم که هرگاه سوالی یا مشکلی برای بنده پیش می‌آمد مجدانه به رفع آن می‌پرداختند.

مجتبی کاشانی

فروردین ۱۴۰۰

تعهد نامه

اینجانب **مجتبی کاشانی** دانشجوی دکتری رشته **آمار علوم ریاضی** دانشگاه شاهرود، نویسنده پایان نامه با عنوان **برآوردگرهای بهبودیافته در مدل های رگرسیونی فازی**، تحت راهنمایی **محمد آرشی و محمدرضا ربیعی** متعهد می شوم:

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش های دیگر پژوهش گران، به مرجع مورد استفاده استناد شده است.
- مطالب این پایان نامه، تا کنون توسط خود، یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ جا ارایه نشده است.
- حقوق معنوی این اثر، به دانشگاه صنعتی شاهرود تعلق دارد، و مقالات مستخرج با نام “دانشگاه صنعتی شاهرود” یا “Shahrood University of Technology” به چاپ خواهد رسید.
- حقوق معنوی تمام افرادی که در به دست آوردن نتایج اصلی پایان نامه تاثیرگذار بوده اند، در مقالات مستخرج از پایان نامه رعایت می گردد.
- در تمام مراحل انجام این پایان نامه، در مواردی که از موجود زنده (یا بافت های آنها) استفاده شده است، ضوابط و اصول اخلاقی رعایت شده است.
- در تمام مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته (یا استفاده شده است)، اصل رازداری و اصول اخلاق انسانی رعایت شده است.

مجتبی کاشانی

فروردین ۱۴۰۰

مالکیت نتایج و حق نشر

- تمام حقوق معنوی این اثر و محصولات آن (مقالات مستخرج، کتاب، برنامه های رایانه ای، نرم افزارها و تجهیزات ساخته شده) متعلق به دانشگاه صنعتی شاهرود می باشد. این مطلب باید به نحو مقتضی، در تولیدات علمی مربوطه ذکر شود.
- استفاده از اطلاعات و نتایج موجود در این پایان نامه بدون ذکر منبع مجاز نمی باشد.

چکیده

در دیدگاه‌های موجود تحلیل رگرسیون فازی، بهبود برآوردگرهای پارامترهای مدل تنها از منظر تصحیح مترهای موجود و برطرف کردن ایرادات فاصله‌ها و مترهای پیشین بوده است. در حالی که در رگرسیون کلاسیک (غیرفازی) بهبود برآوردگرها به منظور دستیابی به بهینگی معیارهای مقایسه و نیکویی برازش از طریق دستکاری و تعمیم خود برآوردگرها نیز صورت می‌گیرد.

بنابراین در این پژوهش برآنیم تا از طریق بهبود برآوردگرها در برخی مدل‌های رگرسیون فازی، و نه متر مورد نظر، برخی معیارهای بهینگی فازی از قبیل شباهت فازی و خطای پیشگویی فازی را بهبود بخشیم. در این راستا استفاده از برآوردگرهای انقباضی و روش‌های باز نمونه‌گیری را در مدل‌های رگرسیون فازی پیشنهاد کرده و نشان می‌دهیم چطور با استفاده از برآوردگر انقباضی فازی و یا باز نمونه‌گیری از داده‌های فازی می‌توان معیارهای بهینگی فازی را بهبود بخشید. همچنین برآوردگرهای جریمه شده فازی را معرفی کرده و دیدگاه انتخاب متغیر فازی را در رگرسیون فازی مورد بررسی قرار دهیم.

کلمات کلیدی: رگرسیون فازی، برآوردگر انقباضی استاین فازی، روش جک نایف بعد از بوت استرپ فازی، برآوردگر جریمه شده.

لیست مقالات مستخرج از پایان نامه

1. Kashani M., Arashi M. and Rabiei M. R. (2021), "Resampling in Fuzzy Regression using by Jackknife-after-Bootstrap: (JB)" **International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems (IJUFKS)**, Accepted on Feb 02, 2021.
 2. Kashani M., Arashi M. and Rabiei M. R. (2021), "An integrated shrinkage strategy for improving efficiency in fuzzy regression modeling" **Soft Computing**, Accepted on Jan 15, 2021.
 3. Kashani M., Arashi M., Rabiei M. R. and D'Urso P. ,and De Giovanni L. (2021), "A Fuzzy Penalized Regression Model with Variable Selection" **Expert Systems With Applications**, 175, pp. 114696-117412.
 4. Kashani M., Arashi M. and Rabiei M. R. (2018). "Jackknife-after-Bootstrap in Fuzzy Regression Modeling" **6th Iranian Joint Congress on Fuzzy and Intelligent Systems**. Shahid Bahonar University of Kerman, Kerman, Iran. 28 Feb-2 Mar.
 5. Kashani M., Arashi M. and Rabiei M. R. (2018). "JB Estimation approach applied to Fuzzy Regression" **8th Seminar on Fuzzy Statistics and Probability**. Department of Statistics, Ferdowsi University of Mashhad, Mashhad, Iran. 9-10 May.
۶. کاشانی م. آرشی م. و ربیعی م. ر. ۱۳۹۸. "یک متر جدید در بررسی فاصله اعداد فازی" **نهمین سمینار ملی آمار و احتمال فازی** دانشگاه مازندران، بابلسر، ایران. ۱۲-۱۱ اردیبهشت ماه.

فهرست مطالب

ش	فهرست تصاویر
ث	فهرست جداول
ذ	پیشگفتار
۱	۱ مدل‌ها و دیدگاه‌های مختلف در رگرسیون فازی
۱	۱.۱ مقدمه
۲	۲.۱ رگرسیون فازی : دیدگاه‌ها و مدل‌ها
۲	۱.۲.۱ انواع مدل‌ها بر اساس ورودی و خروجی
۲	۳.۱ مدل‌های امکانی
۱۶	۴.۱ مدل‌های کمترین فاصله
۲۴	۵.۱ مدل‌های ترکیبی
۲۹	۶.۱ مدل‌های ابتکاری
۳۵	۲ برآورد ضرایب مدل رگرسیون فازی به روش جک نایف بعد از بوت‌استرپ
۳۵	۱.۲ مقدمه
۳۶	۲.۲ مثال‌های انگیزشی
۴۰	۳.۲ بوت‌استرپ در رگرسیون فازی
۴۲	۴.۲ باز نمونه‌گیری در رگرسیون فازی از طریق جک نایف بعد از بوت‌استرپ
۴۴	۱.۴.۲ معیارهای ارزیابی مدل
۴۵	۲.۴.۲ مثال ساختگی ساکاو و یانو (۱۹۹۲)
۴۹	۳.۴.۲ مثال زمان پاسخگویی
۵۵	۴.۴.۲ نتیجه‌گیری
۵۷	۳ مدل‌های رگرسیونی انقباضی فازی
۵۷	۱.۳ مقدمه

۵۸	۱.۱.۳	مدل انقباضی نوع استاین
۶۰	۲.۳	آزمون فرضیه فازی
۶۲	۳.۳	مطالعات عددی
۶۲	۱.۳.۳	مطالعه شبیه سازی
۶۵	۲.۳.۳	تحلیل داده‌های واقعی
۷۵	۴.۳	نتیجه‌گیری
۷۷	۴	مدل‌های رگرسیونی جریمه شده فازی
۷۷	۱.۴	مدل‌های جریمه‌ای
۸۰	۲.۴	مدل رگرسیون جریمه شده فازی
۸۴	۳.۴	معیارهای بهینگی مدل
۸۵	۱.۳.۴	معیار نایقینی
۸۶	۲.۳.۴	فاصله اطمینان t - بوت استری
۸۷	۴.۴	نمایش‌های عددی
۸۸	۱.۴.۴	مطالعه شبیه‌سازی
۹۲	۲.۴.۴	بخش مثال‌های واقعی
۱۰۰	۵.۴	نتیجه‌گیری
۱۰۱	۱.۵.۴	چشم‌اندازهای احتمالی
۱۰۳	آ	تعاریف و مفاهیم اولیه
۱۰۳	۱.آ	مقدمه
۱۰۴	۲.آ	اصل توسیع و اعداد فازی
۱۰۵	۱.۲.آ	برخی فواصل و مترهای فازی
۱۰۶	۲.۲.آ	اعداد فازی
۱۰۸	۳.۲.آ	اعداد فازی LR
۱۰۹	۴.۲.آ	عمل‌های اصلی در اعداد فازی
۱۱۰	۵.۲.آ	حساب اعداد فازی LR
۱۱۰	۳.آ	شاخص‌های نیکویی برازش
۱۱۱	۴.آ	نمودار تیلور
۱۱۳	ب	تعاریف و مفاهیم
۱۱۳	۱.ب	برآوردگرهای بهبودیافته
۱۱۴	۱.۱.ب	برآوردگر انقباضی
۱۱۴	۲.۱.ب	برآوردگر استاین
۱۱۵	۳.۱.ب	تعبیر هندسی برآوردگر نوع استاین

ب.۴.۱	بررسی رفتار برآوردهای انقباضی در توزیع نرمال	۱۱۸
ب.۲	روش‌های باز نمونه‌گیری	۱۲۱
ب.۱.۲	روش جک‌نایف	۱۲۱
ب.۲.۲	روش بوت‌استرپ	۱۲۲
ب.۳.۲	روش جک‌نایف بعد از بوت‌استرپ (JB)	۱۲۲

فهرست تصاویر

۳	درخت رشد رگرسیون فازی در رویکرد امکانی	۱.۱
۴	شرط پوششی مسئله مینیمم	۲.۱
۵	شرط پوششی مسئله ماکزیمم	۳.۱
۶	شرط پوششی مسئله اتصال	۴.۱
۱۶	درخت رشد رگرسیون فازی در رویکرد کمترین فاصله	۵.۱
۲۴	درخت رشد رگرسیون فازی در رویکرد ترکیبی	۶.۱
۲۹	درخت رشد رگرسیون فازی در رویکرد ابتکاری	۷.۱
	۱.۲ مقایسه شهودی مدل های آزموده با مدل های JB متناظر آنها برای مجموعه داده های سوم. (آ) مقدار MPE (ب) مقدار SIM.	۴۷
	۲.۲ مقایسه شهودی فواصل اطمینان بین بوت استرپ و JB (آ) پهنای ضرایب (ب) مرکز ضرایب.	۴۹
	۳.۲ نمودار حبابی با و بدون مشاهده سوم به عنوان نقطه دورافتاده (آ) با نقطه دور افتاده (ب) بدون نقطه دور افتاده.	۵۰
	۴.۲ نمودار فاصله کوک در تشخیص نقطه دورافتاده.	۵۰
	۵.۲ تصویر پاسخ برآورد شده	۵۲
	۶.۲ تصویر عملکرد روش JB در مواجهه با نقطه دورافتاده: (آ) MPE (ب) SIM.	۵۳
	۷.۲ مقایسه فواصل اطمینان بدست آمده در دو روش بوت استرپ و JB (آ) پهنای ضرایب (ب) مرکز ضرایب	۵۴
	۱.۳ مقایسه شهودی سه مدل بر اساس نتایج داده های شبیه سازی: (آ) مقادیر MPE (ب) مقادیر SIM.	۶۳
	۲.۳ تحلیل شهودی مسیر p -مقدار.	۶۴
	۳.۳ مقایسه شهودی مدل انقباضی و مدل پایه بر اساس معیارهای ارزیابی در مثال اول: (آ) مقدار MPE (ب) مقدار SIM.	۶۶
	۴.۳ تحلیل شهودی مسیر معیار p -مقدار.	۶۸

۵.۳	مقایسه شهودی و استنباطی مدل پایه و مدل های انقباضی در مقادیر مختلف k ، در مثال اول. (آ) مقدار مرکز (ب) مقدار پهنا.	۶۸
۶.۳	مقایسه شهودی مدل پایه و انقباضی تحت معیارهای ارزیابی. (آ) مقدار MPE (ب) مقدار SIM.	۷۰
۷.۳	تحلیل شهودی مسیر معیار p - مقدار.	۷۰
۸.۳	تصویر پاسخ برآورد شده.	۷۲
۹.۳	مقایسه شهودی و استنباطی مدل های پایه و انقباضی در مقادیر مختلف k در مثال دوم:	۷۲
۱۰.۳	مقایسه شهودی مدل پایه و انقباضی تحت معیارهای ارزیابی (آ) مقدار MPE (ب) مقدار SIM.	۷۴
۱۱.۳	مقایسه شهودی و استنباطی مدل های انقباضی در مقادیر مختلف k با مدل مبنا در مثال سوم: (آ) مقدار MPE (ب) مقدار SIM.	۷۵
۱.۴	مقایسه شهودی - استنباطی مدل های جریمه ای با مدل های رقیب. (آ) مقدار مرکز (ب) مقدار پهنا.	۹۵
۱.آ	نمودار تابع نیم پیوسته پایینی	۱۰۷
۲.آ	نمودار تابع نیم پیوسته بالایی	۱۰۷
۳.آ	نمودارهای اعداد فازی تقریباً صفر.	۱۰۸
۴.آ	عدد فازی مثلثی تقریباً صفر.	۱۰۸
۵.آ	عدد فازی نرمال تقریباً صفر.	۱۰۹
۶.آ	عدد فازی سهموی تقریباً صفر.	۱۰۹
۱.ب	توجیه هندسی برآوردگر انقباضی استاین	۱۱۵
۲.ب	نمودار میانگین ها به روش انقباضی	۱۱۹

فهرست جداول

۱.۲	مجموعه داده‌های زمان تشخیص خدمه اتاق کنترل نیروگاه هسته ای بر
۳۶	یک واقعه غیرعادی [۸۲]
۳۷	مجموعه داده دوم.
۳۷	مجموعه داده‌های ساختگی ساکاو و یانو [۱۱۶]
۳۸	مجموعه داده‌های قیمت خانه‌های شانگهای [۱۵۱]
۴۲	معیارهای ارزیابی برای مدل بوت استرپی، در مجموعه داده‌های سوم ۳.۲.
۴۲	معیارهای ارزیابی مدل قبل و بعد از بکارگیری روش JB برای مجموعه
۴۵	داده‌های سوم.
۷.۲	ضرایب برآورد شده برای مدل‌های مبنا و مدل‌های JB متناظر آن‌ها در
۴۶	مجموعه داده‌های سوم.
۴۷	فواصل اطمینان بوت استرپی برای پارامترها.
۵۱	نتایج معیارهای ارزیابی GOF قبل و بعد از حذف نقطه دورافتاده.
۱۰.۲	فواصل اطمینان t -بوت استرپی. (مقادیر داخل پرانتز طول فاصله اطمینان
۵۳	است)
۱.۳	مقادیر GOF برای ارزیابی میانگین ۳۰ مجموعه داده شبیه‌سازی شده برای
۶۲	روش انقباضی در مقابل مدل‌های دیگر.
۶۳	نتایج ارزیابی روش انقباضی در ۱۸ امین مجموعه داده شبیه‌سازی شده.
۳.۳	نتایج آزمون فرضیه در داده‌های شبیه‌سازی شده ۱۸ ام برای $\tilde{\beta}_{j\alpha}$ همراه با
۶۴	T_α و معیار p -مقدار.
۴.۳	مقایسه مدل انقباضی و مدل پایه بر اساس معیارهای ارزیابی GOF در مثال
۶۵	اول.
۵.۳	آزمون فرض هر $\tilde{\beta}_{j\alpha}$ همراه با T_α و معیار p -مقدار.
۶.۳	مدل‌های برازش شده مبنا، بهینه انقباضی و IFS برای مثال اول.
۶۹	مقادیر GOF برای ارزیابی مدل‌های پایه انقباضی در مثال دوم.
۷۰	آزمون فرض هر ضریب $\tilde{\beta}_{j\alpha}$ همراه با T_α و معیار p -مقدار.

۷۱	۹.۳	مدل‌های برازش شده برای هر ثابت انقباضی بهینه در مثال دوم.
۷۳	۱۰.۳	ارزیابی مدل‌های پایه و انقباضی تحت معیارهای شش گانه ارزیابی
۷۴	۱۱.۳	مدل‌های برازش شده توسط روش پایه، انقباضی نوع استاین و مدل IFS.
	۱.۴	مقایسه مدل‌های جریمه با مدل‌های رقیب بر اساس میانگین معیارهای ارزیابی
۹۰		بهینگی در مجموعه داده‌های شبیه‌سازی
۹۱	۲.۴	نتایج مدل‌های جریمه ی ۸ ام، ۱۵ ام، ۱۱ ام، به ترتیب در سناریوهای اول تا سوم.
۹۳	۳.۴	معیارهای ارزیابی مدل جهت مقایسه مدل‌های جریمه و مدل‌های رقیب
۹۳	۴.۴	ضرایب مدل‌های جریمه‌ای برای داده‌های مزه پنیر
	۵.۴	فاصله اطمینان t - بوت استری برای پارامترها در مدل‌های مختلف (اعداد داخل
۹۴		پرانتز طول فاصله اطمینان است).
۹۵	۶.۴	مقدار se برای ضرایب مدل‌ها در داده‌های مزه پنیر
۹۷	۷.۴	نتایج ارزیابی مدل‌های جریمه شده و مدل‌های رقیب به همراه ضرایب مدل.
۹۸	۸.۴	مقدار se برای ضرایب متغیرهای داده‌های قیمت خانه‌های شانگهای.
	۹.۴	فاصله اطمینان t - بوت استری برای پارامترها در مدل‌های مختلف (اعداد داخل
۹۹		پرانتز طول فاصله اطمینان است).
۱۱۷	۱.ب	نتایج بازی بیس بال مربوط به مثال ب.۱.۱.

پیشگفتار

انسان از دیرباز، به واسطه ذات کنجکاو خود، به دنبال کشف حقایق و مسائل طبیعت پیرامون خود بوده است. او با نگاه کردن به رفتارهای دیگر موجودات سعی در یادگیری و علاوه بر آن رسیدن به ماهیت روابط موجود در این پدیده‌ها بوده است و همین امر او را نسبت به کشف روابط و ساختارهای آن‌ها ترغیب نموده است. این مسئله سرآغاز ظهور علوم مختلف بوده است. از سوی دیگر، علاوه بر سعی در کشف رازهای علمی، مهمترین مسئله پیش‌بینی رخداد‌های مختلف بوده است. زیاده‌گویی نیست اگر بگوییم که مهمترین بخش در هر علمی پیشگویی اتفاقات ممکن در آن پدیده است. در این بین شاخه مهمی از علم، به نام علم آمار خودنمایی می‌کند که برای پوشش پدیده‌های مختلف همواره یکی از مبانی مهم تحقیق علمی است؛ زیرا با نگاه به پدیده‌های طبیعی و الهام از آن‌ها شاخه‌های مختلف مدل‌بندی و پیش‌بینی از جمله رگرسیون را ارائه داده است. در این مسیر، برای ایجاد ساختارهایی علمی، تعریف و بیان پدیده‌ها و افزایش دقت خود در برآورد پارامترهای موثر بر مدل‌بندی و در نهایت پیش‌بینی رفتار آن‌ها، فرمول‌ها و قواعدی را ابداع نمود. از جمله شرایط اعمال شده بر تولید یک برآورد بهینه برای پارامترهای مجهول می‌توان به شرط ناریبی پرداخت. این شرط از آنجایی که در مباحث نظری نمی‌توان به محاسبه دقیق پارامتر از طریق کمترین توان دوم خطای برآورد دست یافت، ایجاد گردید. اما همواره ایجاد چنین شرایط نامنعطف، به نوعی حذف تعدادی از برآوردها را به دنبال دارد که ممکن است در موارد مختلف نسبت به آن‌ها عملکرد بهتری داشته باشند. اما، هرگاه انسان در رفع مشکلات ایجاد شده در کارهای خود نگاه عمیق‌تری به پدیده‌های طبیعی داشت و سعی کرد تا از رفتار آن‌ها در مباحث خود تقلید کند، به نتایج شگفت‌آوری دست یافت. این موانع و مشکلات می‌تواند شامل تعداد کم نمونه (چون در موارد زیادی برای صفات مورد بررسی، قابلیت دریافت نمونه کافی وجود ندارد)، وجود همبستگی بین متغیرهای پیشگو که به همخطی نیز معروف است و از موارد آن می‌توان افزایش تعداد متغیرها و یا وجود همبستگی ذاتی بین متغیرهای موثر نام برد، در مدل‌بندی رگرسیون باشد. از این رو و برای رفع آن‌ها روش‌هایی ابداع گردید که بکارگیری از آن‌ها بهبودی برآوردهای پارامترها را به دنبال داشت. از جمله این موارد می‌توان از روش‌های بازنمونه‌گیری شامل جک نایف و بوت استرپ و جک نایف بعد از بوت استرپ و روش‌های بهبود برآوردها شامل برآوردگرهای انقباضی و برآوردگرهای جریمه‌ای که معمولاً برآوردگرهایی اریب را تولید می‌کنند،

نام برد. در واقع این روش‌ها دیدگاه‌های متفاوتی دارند. روش‌های باز نمونه‌گیری از طریق دستیابی به توزیع تجربی و در روش‌های انقباضی و جریمه‌ای از طریق منعطف کردن برآوردگر نااریب مانند شگردهای انعطاف‌پذیری که طبیعت برای رفع موانع روبروی خود ابداع می‌کند و دست کاری در آن با تغییر در مقدار پارامتر، معروف به پارامتر تنظیم، به فرآیند بهینه‌سازی مدل در خصوص تطابق بهتر با مشاهدات واقعی و در نتیجه پیش‌بینی‌هایی دقیق‌تر کمک می‌کنند.

از سوی دیگر، بسیاری از پدیده‌ها و صفات ماهیتی مبهم دارند و بکارگیری و اندازه‌گیری آن‌ها از طریق مترهای دقیق، نادرست و در مواردی باعث خطا در برآورد و ارائه پیش‌بینی‌های نادرست می‌شود. بنابراین، در اینجا استفاده از داده‌های فازی، منطقی و درست به نظر می‌رسد. کما اینکه، در مواردی مثلاً در مدل‌بندی رگرسیونی ممکن است شرایط توزیعی نیز برقرار نباشد و این مسئله نیز ما را مجاب به استفاده از روش‌هایی مانند روش‌های فازی می‌کند که فارغ از برقراری شرایط توزیعی است.

در این رساله با توجه به مطالب عنوان شده در فصل اول یک بازنگری کلی در خصوص روند شروع و پیشرفت رگرسیون فازی انجام شده است. در اکثر مدل‌های پیشنهادی مبنای بهینه‌سازی برآورد پارامترها برای برازش بهتر مدل و پیش‌بینی بهتر آن‌ها، تغییر در مترهای به کار رفته و یا تغییر شرایط بررسی مدل‌ها است. بنابراین در فصل‌های بعدی بر آن شدیم تا با دیدگاه جدید، بدون تغییر در مترها و یا شرایط اولیه مسئله، از طریق سه استراتژی به بهبود برآوردهای مدل و برازش آن بپردازیم. در فصل دوم از طریق روش جک نایف بعد از بوت استرپ علاوه بر رفع مشکل بوت استرپ در تصادفی بودن انحراف استاندارد تولید شده برای برآوردها، به بهبود مدل از این روش و بررسی استوار بودن این روش طی مثال‌هایی پرداخته‌ایم. در ادامه و در فصل سوم بکارگیری رویکرد انقباضی با عنوان برآورد انقباضی نوع استاین فازی و مدل انقباضی یکپارچه (IFS)، بهبودی برآوردها و معیارهای ارزیابی مدل را در مقابل مدل‌های دیگر طی مثال‌های متنوع واقعی و شبیه‌سازی شده نشان داده‌ایم. در فصل چهارم نیز با بکارگیری مدل‌های جریمه شده علاوه بر بهبود مدل‌ها در مقابل مدل‌های رقیب، انتخاب متغیر خودکار را هم در مثال‌هایی با داده‌های شبیه‌سازی شده و واقعی با استفاده از مدل‌های جریمه شده نشان داده‌ایم.

فصل ۱

مدل‌ها و دیدگاه‌های مختلف در رگرسیون فازی

۱.۱ مقدمه

در ریاضیات مدرن، نظریه مجموعه‌های معمولی به عنوان پایه و شالوده اصلی آن است. اما همواره با ویژگی‌های مطلق روبرو نیستیم، مانند کوتاه بودن، کافی بودن و ... همچنین در علوم رفتاری، روان‌شناسی و دیگر حوزه‌ها با مفاهیم و تعاریف نادقیقی سروکار داریم که در قالب نظریه مجموعه‌های معمولی جایگاهی ندارند.

بدین شکل قالب جدیدی برای پوشش این مفاهیم نادقیق احساس می‌گردد. نظریه مجموعه‌های فازی این امر را عهده‌دار شده و پوشش مناسبی به این مجموعه‌ها ارائه می‌دهد. معرفی و تشریح این نظریه در سال ۱۹۶۵ میلادی (۱۳۴۴ ه.ش.) توسط پروفیسور لطفی عسگرزاده [۱۴۹] دانشمند ایرانی تبار عرشه شد. این نظریه از زمان ارائه تاکنون، گسترش و تعمیم زیادی یافته و کاربردهای گوناگونی در زمینه‌های مختلف پیدا کرده است. به طور خلاصه، نظریه‌ی مجموعه‌های فازی نظریه‌ای در شرایط عدم اطمینان است. این نظریه قادر است بسیاری از مفاهیم و متغیرها و سیستم‌هایی را که نادقیق و مبهم هستند، همان‌گونه که در عالم واقع اکثراً چنین است، صورت‌بندی ریاضی بخشیده و زمینه را برای استدلال، استنتاج، کنترل و تصمیم‌گیری در شرایط عدم اطمینان فراهم آورد (برای اطلاعات بیشتر به

پیوست آ مراجعه کنید).

۲.۱ رگرسیون فازی : دیدگاه‌ها و مدل‌ها

یکی از کاربردهای روش‌های فازی، بکارگیری آن‌ها در ارائه مدل‌های رگرسیونی است؛ چرا که بسیاری از متغیرها در واقعیت، ماهیتی مبهم دارند. به علاوه ممکن است همه شرایط اولیه در رگرسیون کلاسیک خصوصا شرایط توزیعی در مسئله برقرار نباشد. در این قسمت مرور اجمالی بر روند توسعه رگرسیون فازی از زوایای مختلف خواهیم داشت و این روند را هم به صورت نموداری و هم با ارائه گزارشی مختصر از دلایل آن، در چهار دیدگاه‌ها امکانی، کمترین فاصله، ترکیبی و ابتکاری بیان خواهیم نمود.

اکنون برای جمع‌بندی این دیدگاه‌ها، در ابتدای هر بخش درخت رشد متناظر با هر دیدگاه ارائه شده است تا شمای کلی از آن‌ها روند تغییرات و مدل‌های مختلف را داشته باشیم.

۱.۲.۱ انواع مدل‌ها بر اساس ورودی و خروجی

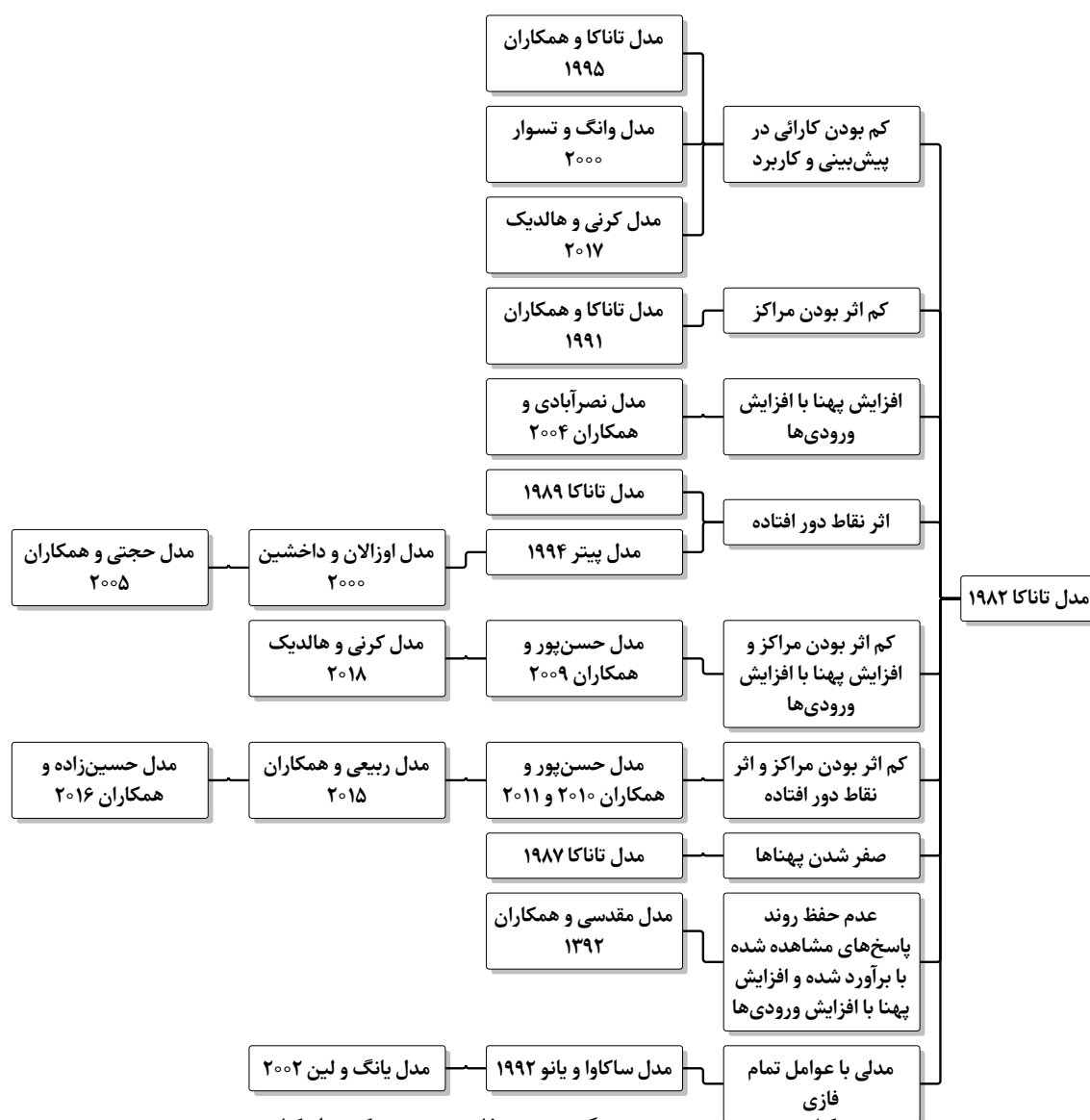
انواع مدل‌های ممکن در رگرسیون فازی به شرح زیر ارائه می‌گردد:

۱. ورودی غیرفازی، خروجی و ضرایب فازی
۲. ورودی و خروجی فازی، ضرایب غیرفازی
۳. ورودی و خروجی غیرفازی، ضرایب فازی
۴. ورودی، خروجی و ضرایب همگی فازی

۳.۱ مدل‌های امکانی

در این روش‌ها تحت شرایطی و بر اساس یک مدل برنامه‌ریزی خطی یا غیرخطی، ضرایب لازم برای می‌نیم یا ماکزیمم کردن یک تابع هدف به دست خواهد آمد. قبل از بررسی انواع مدل‌ها روند آن‌ها را در درخت رشد مدل‌های امکانی به شرح زیر مشاهده می‌کنیم.

درخت رشد رگرسیون فازی در رویکرد امکانی



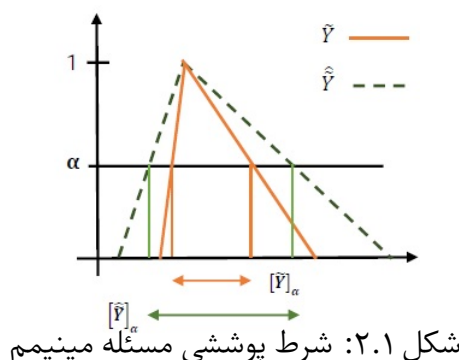
شکل ۱.۱: درخت رشد رگرسیون فازی در رویکرد امکانی

مدل تاناکا و همکاران (۱۹۸۲)

در این مدل ماتریس مشاهدات مستقل، اعدادی دقیق به صورت $X = [x_{ij}]_{n \times p}$ است، که در آن $i = 1, 2, \dots, n$ و $j = 0, 1, 2, \dots, p$ و $x_{i0} = 1$ ، همچنین $\tilde{A} = \{\tilde{A}_0, \tilde{A}_1, \dots, \tilde{A}_p\}$ مجموعه ضرایب مدل است و به صورت اعداد فازی مثلثی $\tilde{A}_j = (a_j, s_j)_T$ فرض شده‌اند. این مقاله در دو حالت بحث شده است. در حالت اول متغیر پاسخ مقداری دقیق و در حالت دوم یک عدد فازی مثلثی به صورت $\tilde{Y}_i = (y_i, e_i)_T$ است. مدل رگرسیونی را به شرح زیر در نظر بگیرید:

$$\tilde{Y} = \tilde{A}_0 + \tilde{A}_1 x_1 + \dots + \tilde{A}_p x_p \quad (1.1)$$

تحت شرایط زیر ضرایب مدل به دست می‌آید:



شکل ۲.۱: شرط پوششی مسئله مینیمم

(الف) برای مدل با متغیر پاسخ دقیق تابع عضویت پاسخ برآورد شده دارای حداقل α درجه عضویت باشد، به عبارت دیگر

$$\forall y_i, \quad \hat{Y}_i(y_i) \geq \alpha \quad (2.1)$$

(ب) برای مدل با متغیر پاسخ فازی می‌بایست α -برش پاسخ برآورد شده شامل α -برش پاسخ مشاهده شده باشد (۲.۱) یعنی

$$[\tilde{Y}_i]_{\alpha} \subseteq [\hat{Y}_i]_{\alpha} \quad (3.1)$$

(ج) ضرایب به گونه‌ای بدست خواهند آمد که تابع هدف، می‌نیمم ابهام کل مدل باشد یعنی،

$$\min Z = \sum_{j=0}^p e_j \quad (4.1)$$

در این روش مشکلات متعددی وجود داشت که روش‌ها و پیشنهادهای متعددی در جهت رفع و بهبود آن‌ها ارائه گردید که به شرح زیر است،

(الف) صفر شدن پهنایها (عدم سنخیت با مفهوم فازی) به واسطه وجود قید $s_j \geq 0$ ، جواب‌های حاصل برای می‌نیمم کردن تابع هدف می‌تواند صفر شود که این مسئله به خصوص در مواردی که پاسخ فازی است می‌تواند ما را از واقعیت دور کند [۱۲۸].

(ب) عدم توجه به نقاط مرکزی و کاهش راندمان پیش بینی در تابع عضویت یک عدد فازی، مراکز دارای بالاترین درجه عضویت هستند که به این مسئله در روش فوق اشاره نشده است [۱۲۰]. این مسئله باعث کاهش بازده پیش بینی این روش نسبت به روش‌های غیرفازی می‌شود [۱۳۷].

(ج) حساسیت نسبت به نقاط دور افتاده:

از آنجایی که در روش مذکور، برآورد پهنای متغیر پاسخ از مجموع پهنای ضرایب برآورد شده مدل حاصل می‌گردد لذا به واسطه‌ی وجود شرط پوششی (در شرایط مدل) برای پوشش تمام نقاط، در صورت وجود نقطه‌ی دور افتاده نیاز به افزایش پهنای برآوردی داریم که در واقع یک بیش برآورد خواهیم داشت [۸۱، ۱۱۶، ۱۲۹].

(د) افزایش پهنای پاسخ‌های برآورد شده با افزایش تعداد متغیر مستقل:

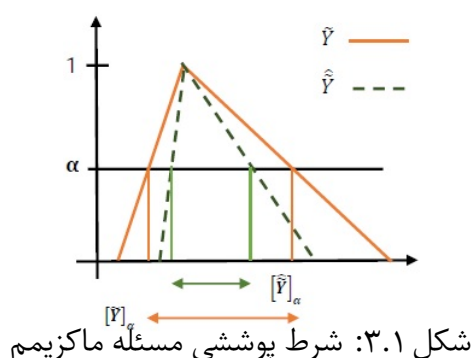
با توجه به تابع هدف مدل زمانی که ضرایب مدل اعداد فازی باشند، با افزایش تعداد متغیرهای مستقل، پهنای پاسخ‌های برآورد شده افزایش می‌یابد. این مسئله در زمانی که پهنای پاسخ‌های مشاهده شده کاهش یابد و یا ثابت باشد مسئله‌ساز خواهد بود [۷۷].

مدل تاناکا (۱۹۸۷)

علیرغم سادگی محاسبات مدل (۱.۱)، تاناکا [۱۲۸] برای رفع مشکل صفر شدن پهنای تابع هدف را از ۴.۱ به $\min Z = \sum_{j=0}^p s_j |x_{ij}|$ تغییر داد. در واقع بجای ابهام کل پاسخ، ابهام کل برآورد شده مورد هدف قرار گرفت که تا حد زیادی به رفع این مسئله کمک کرد. علاوه بر آن، مدل براساس تابع هدف ماکزیمم ارائه گردید که شرط پوششی آن‌ها معکوس است:

$$[\hat{\bar{Y}}_i]_{\alpha} \subseteq [\tilde{Y}_i]_{\alpha} \quad (5.1)$$

این مسئله عکس مسئله می‌نیم است، در نتیجه تابع هدف به $\max Z = \sum_{j=0}^p s_j x_{ij}$ تغییر می‌یابد.

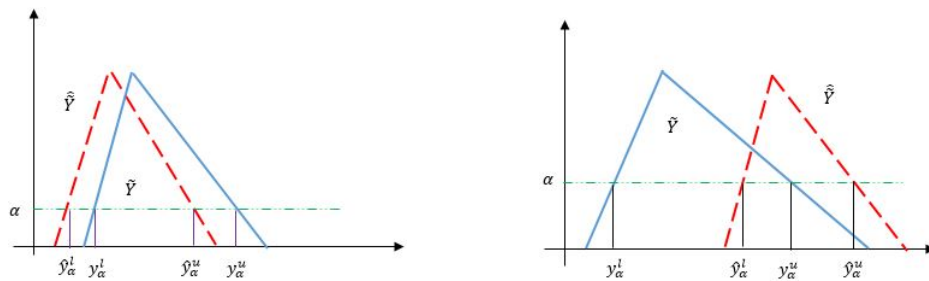


مزیت مسئله ماکزیمم نسبت به مسئله می‌نیم آن است که مقدار تابع هدف در مسئله ماکزیمم کمتر یا مساوی مقدار تابع هدف مسئله می‌نیم است.

مدل تاناکا (۱۹۸۹)

پس از رفع مشکل صفر شدن پهناها، تاناکا و همکاران در [۱۲۹]، برای رفع مشکل اثر نقاط دور افتاده در افزایش پهناهای ضریب برآورد شده بر روی شرط پوشش بین تابع عضویت پاسخ‌های مشاهده شده و برآورد شده نسبت به هم، کار کردند و مسئله اتصالی را ارائه کردند که به صورت زیر نشان داده می‌شود.

مسئله اتصال (قطع کردن): در این حالت تنها کافی است تحت هر α - برش دلخواه حد بالای پاسخ برآورد شده از حد پایین پاسخ مشاهده شده بزرگتر ($\hat{y}_\alpha^u > y_\alpha^l$) و حد پایین پاسخ برآورد شده از حد بالای پاسخ مشاهده شده کمتر باشد ($\hat{y}_\alpha^l > y_\alpha^u$) (شکل ۴.۱).



شکل ۴.۱: شرط پوششی مسئله اتصال

در این مقاله تاناکا نشان داد که مقدار بهینه تابع هدف در مسئله اتصال و ماکزیمم کمتر از مقدار بهینه در مسئله می‌نیمم است. همچنین مزیت روش اتصال نسبت به مسئله می‌نیمم آن است که در حضور نقاط دور افتاده از افزایش پهناهای نامتعارف تا حدی جلوگیری می‌نماید.

مدل ساکاو و یانو (۱۹۹۲)

در مدل‌های قبلی مقادیر متغیرهای ورودی همگی دقیق فرض شده بودند از این‌رو ساکاو و یانو [۱۱۶]، با در نظر گرفتن مدل ۶.۱، به تعمیم مدل‌های قبل پرداختند.

$$\tilde{y}_i = \tilde{A} \otimes \tilde{X}_i, \quad i = 1, 2, \dots, n \quad (6.1)$$

که در آن $\tilde{y}_i = (y_i, e_i)_L$ ، $\tilde{A} = (a_j, s_j)_L$ ، $\tilde{A} = (\tilde{A}_0, \tilde{A}_1, \dots, \tilde{A}_p)$ ، $\tilde{X}_i = (1, \tilde{X}_{i1}, \dots, \tilde{X}_{ip})$ و $\tilde{X}_{ij} = (x_{ij}, d_{ij})$. بنابراین ورودی، خروجی و ضرایب مدل عدد فازی متقارن هستند و $\otimes$ عملگر ضرب بر اساس اصل گسترش است. علاوه بر آن فرض بر آن است که ورودی‌ها اعداد فازی مثبت باشند، همچنین به واسطه آنکه عمل ضرب دو عدد فازی LR ، تقریباً LR است، لذا استفاده از اصل گسترش دشوار خواهد بود. از این‌رو برای حل مشکل از روش [۹۹] استفاده گردید که در آن نشان داده است که تحت شرایطی $(A \times B)_\alpha = (A_\alpha \times B_\alpha)$. بنابراین

می‌شود که در آن تحت α -برش دلخواه داریم:

$$\begin{aligned} y_{i\alpha}^L &= \sum_{j=\circ}^p \left(\min \left((a_j - L^{-1}(\alpha)s_j)(x_{ij} - L^{-1}(\alpha)d_{ij}), (a_j - L^{-1}(\alpha)s_j)(x_{ij} + L^{-1}(\alpha)d_{ij}) \right) \right) \\ y_{i\alpha}^R &= \sum_{j=\circ}^p \left(\max \left((a_j + L^{-1}(\alpha)s_j)(x_{ij} + L^{-1}(\alpha)d_{ij}), (a_j + L^{-1}(\alpha)s_j)(x_{ij} - L^{-1}(\alpha)d_{ij}) \right) \right) \end{aligned} \quad (7.1)$$

که به ترتیب $y_{i\alpha}^R$ و $y_{i\alpha}^L$ می‌نیمیم و ماکزیمم مقدار عملگر در برش دلخواه α برای $(\tilde{A} \otimes \tilde{X}_i)_\alpha$ است و در آن $L^{-1}(\alpha)$ معکوس تابع عضویت به ازای برش α است. در این مقاله برای مشخص شدن مقادیر (7.1)، سه گروه زیر تعیین گردیده است،

$$a_j - L^{-1}(\alpha)s_j \geq \circ \quad j \in J_1 \quad (8.1)$$

$$a_j - L^{-1}(\alpha)s_j \leq \circ \quad a_j + L^{-1}(\alpha)s_j \geq \circ \quad j \in J_2 \quad (9.1)$$

$$a_j + L^{-1}(\alpha)s_j \leq \circ \quad j \in J_3 \quad (10.1)$$

که در آن،

$$J = J_1 \cup J_2 \cup J_3 = \{\circ, 1, \dots, p\}, \quad J_k \cap J_{k'} = \emptyset, \quad k(\neq k') = 1, 2, 3 \quad (11.1)$$

در هر سه دسته، تابع هدف عبارت است از

$$\min Z = \sum_{i=1}^n (y_{i\circ}^R - y_{i\circ}^L) \quad (12.1)$$

و شرایط ارائه شده برای تابع هدف (12.1) براساس سه نوع پوشش بین $\tilde{y}_i$ و $\hat{\tilde{y}}_i$ لحاظ گردیده است. آن‌ها در نهایت الگوریتمی ارائه کردند که حالات فوق را بررسی و برآوردهای مورد نظر تحت یک برنامه‌ریزی خطی تولید کند.

مدل پیتز (۱۹۹۴)

در سال ۱۹۹۴ پیتز پس از بررسی روش تاناکا [۱۲۸] در [۱۰۳] به این موضوع اشاره کرد که محدودیت‌های روش تاناکا باعث می‌شود تمام پاسخ‌های مشاهده شده در یک فاصله مشخص قرار گیرند که در صورت وجود داده دور افتاده در مشاهدات، این فاصله متاثر از آن، دارای پهنای زیاد برای برآوردها خواهد بود. در این مسیر مدل بکار رفته با مشاهدات مستقل و پاسخ‌های مشاهده شده دقیق است و ضرایب مدل اعداد فازی مثلی فرض شده‌اند. در این مدل فرض می‌شود که حد پایین مقدار پاسخ برآورد شده از مقدار پاسخ مشاهده شده

متناظر آن کمتر و حد بالای آن از مقدار متناظرش در پاسخ مشاهده شده بیشتر باشد. مدل به شرح زیر است،

$$\begin{aligned} \max \quad & \bar{\lambda} = \frac{1}{n} \sum_{i=1}^n \lambda_i \\ \text{s. t.} \quad & (1 - \bar{\lambda})p_0 - \sum_{i=1}^n \sum_{j=0}^p s_j |x_{ij}| \geq 0 \\ & (1 - \lambda_i)p_i + \sum_{j=0}^p a_j x_{ij} + \sum_{j=0}^p s_j |x_{ij}| \geq y_i \\ & (1 - \lambda_i)p_i - \sum_{j=0}^p a_j x_{ij} + \sum_{j=0}^p s_j |x_{ij}| \geq -y_i \\ & s_i \geq 0, \quad a_j \in \mathbb{R}, \quad \lambda \leq 1, \quad x_{i0} = 1 \end{aligned} \quad (13.1)$$

که در آن λ درجه عضویت مربوط به حل مسئله است. p_i و p_0 محدوده تغییرات فاصله‌ای که شامل مشاهدات است و با افزایش p_i و مقدار p_0 کمتر $(p_0, p_i \in \mathbb{N})$ ، حدود کمتر شده و پاسخ‌های برآورد شده سخت‌گیرانه‌تر می‌شود و این مقادیر توسط کاربر انتخاب می‌گردد.

مدل تاناکا و همکاران (۱۹۹۵)

با هدف ارائه کاربردی برای رگرسیون امکانی تاناکا و همکاران [۱۳۱] در سال ۱۹۹۵، براساس توزیع نمایی به تحلیل رگرسیون پرداختند. در واقع هدف، نشان دادن تحلیل و تفسیرپذیری رگرسیون امکانی به خوبی تحلیل آماری بود. بدین معنی که این روش برای تحلیل پدیده‌های اجتماعی و اقتصادی جایگزین مناسبی بجای تحلیل آماری است.

مدل اوزالان و داخشتین (۲۰۰۰)

در مدل اوزالان و داخشتین [۱۰۰] شرط قطع شدن پاسخ‌های برآورد شده و مشاهده شده وجود دارد. علاوه بر آن در [۱۰۳] برای دریافت مقادیر ثابت p_i در مدل، روش عملی ارائه نشده است و مهم‌تر از آن حساسیت به داده دور افتاده تا حدودی برطرف شده است. [۱۰۰] روشی بر مبنای رگرسیون فازی با چند هدف ارائه گردید است. علاوه بر آن، آن‌ها مسئله را در دو گروه بررسی کردند. گروه اول بنام رگرسیون فازی با دو هدف و گروه دوم رگرسیون فازی با سه هدف که هر کدام نیز تحت سه مدل برنامه‌ریزی (براساس قیدها و عوامل موجود در تابع هدف) ارائه گردید. همچنین مزیت دیگر این روش نسبت به روش پیتر [۱۰۳] عدم نیاز به مشخص کردن پارامترهای خاص از سوی محقق است.

مدل وانگ و تسوار (۲۰۰۰)

این مدل که در سال ۲۰۰۰ ارائه گردید [۱۳۷] با هدف تعدیل کردن روش کمترین توان های دوم براساس یک مسئله برنامه‌ریزی تحت روش تاناکا است که از لحاظ پیش‌بینی از مدل تاناکا، و از لحاظ بازده محاسباتی از روش کمترین توان های دوم مرسوم بهتر است. بر این اساس مدل با فرض ورودی دقیق، ضرایب و خروجی فازي $\tilde{y}_i = (y_i, e_i)_L$ در نظر گرفته شد و مدل آن عبارت است از

$$\begin{aligned} \min \quad & Z = \sum_{i=1}^n \left(\sum_{j=0}^p s_j |x_{ij}| - e_i \right)^2 \\ \text{s. t.} \quad & \sum_{j=0}^p a_j x_{ij} + \sum_{j=0}^p s_j |x_{ij}| \gtrsim y_i + e_i \\ & - \sum_{j=0}^p a_j x_{ij} + \sum_{j=0}^p s_j |x_{ij}| \gtrsim -y_i + e_i \\ & s_j \geq 0, \quad a_j \in \mathbb{R}, \quad x_{i0} = 1, \quad \forall i = 1, 2, \dots, n, j = 0, 1, \dots, p \end{aligned} \quad (14.1)$$

که در آن $\gtrsim$ به معنی تقریباً بزرگتر است.

مدل لی و چن (۲۰۰۱)

در سال ۲۰۰۱ لی و چن با فرض مدل با ورودی و خروجی فازي [۸۴] تحت یک برنامه‌ریزی غیرخطی با تابع هدفی که در آن مجموع پهنای ضرایب مدل را می‌نیمم کند ارائه کردند. البته در سال ۲۰۰۳ هونگ و چویی طی یک مثال نقض، نامعتبر بودن این روش را نشان دادند.

مدل شن یانگ و شون لین (۲۰۰۲)

ساکاوا و یانو در [۱۱۶] مدلی را ارائه دادند که در آن تمامی عوامل مدل، فازي بودند و تنها مسئله، مثبت بودن عدد فازي متغیر ورودی بود. لذا شن یانگ و شون لین در [۱۴۴] این شرط را حذف کرده و با کلاسه‌بندی مشاهدات (با فرض ناهمگونی مشاهدات) و ارائه یک وزن به هر کلاسه تحت معیار کمترین توان دوم فازي با دو ساختار فاصله تقریبی دو عدد فازي و فاصله بازه‌ای (آنالیز بازه‌ای) مدل ساکاوا و یانو را تعمیم دادند و این دو ساختار را مقایسه نموده که در آن ساختار فاصله تقریبی عملکرد بهتری از فاصله بازه‌ای داشته است.

مدل چانگ وو (۲۰۰۳)

چانگ وو در [۱۴۰] با فرض فازي بودن تمام عوامل مدل رگرسیون خطی، با استفاده از روش کمترین توان های دوم در رگرسیون معمولی و اینکه برای هر دو عدد فازي $\tilde{A}$ و $\tilde{B}$ عملگر جمع

و ضرب را به صورت زیر، تحت α -برش دلخواه داریم،

$$\begin{aligned} (\tilde{A} \oplus \tilde{B})_{\alpha} &= [\tilde{A}_{\alpha}^L + \tilde{B}_{\alpha}^L, \tilde{A}_{\alpha}^U + \tilde{B}_{\alpha}^U] \\ (\tilde{A} \otimes \tilde{B})_{\alpha} &= [\min \{ \tilde{A}_{\alpha}^L \tilde{B}_{\alpha}^L, \tilde{A}_{\alpha}^U \tilde{B}_{\alpha}^L, \tilde{A}_{\alpha}^L \tilde{B}_{\alpha}^U, \tilde{A}_{\alpha}^U \tilde{B}_{\alpha}^U \}, \max \{ \tilde{A}_{\alpha}^L \tilde{B}_{\alpha}^L, \tilde{A}_{\alpha}^U \tilde{B}_{\alpha}^L, \tilde{A}_{\alpha}^L \tilde{B}_{\alpha}^U, \tilde{A}_{\alpha}^U \tilde{B}_{\alpha}^U \}] \end{aligned}$$

و در نتیجه تحت α -برش دلخواه داریم،

$$(\tilde{y}_i)_{\alpha}^L = \beta_{\alpha}^L + \beta_{\alpha}^L (\tilde{X}_{i1})_{\alpha}^L + \cdots + \beta_{p,\alpha}^L (\tilde{X}_{ip})_{\alpha}^L$$

$$(\tilde{y}_i)_{\alpha}^U = \beta_{\alpha}^U + \beta_{\alpha}^U (\tilde{X}_{i1})_{\alpha}^U + \cdots + \beta_{p,\alpha}^U (\tilde{X}_{ip})_{\alpha}^U$$

برآورد پارامترها را تحت تابع هدف $z = \max \alpha$ بررسی نموده و در قالب چهار مدل برنامه‌ریزی خطی آن را ارائه داد.

مدل نصرآبادی و نصرآبادی (۲۰۰۴)

نصرآبادی و نصرآبادی [۹۶] با اشاره به این مشکل که در مدل‌های قبل [۷۸، ۱۲۸] با افزایش مقادیر متغیر مستقل، پهنای پاسخ‌های برآوردشده افزایش می‌یابد و ممکن است این روند در مشاهدات پاسخ وجود نداشته باشد با ارائه یک عملگر جبری برای ضرب روی اعداد فازی متقارن مدل رگرسیونی فازی زیر را مورد بررسی قرار دادند.

$$\tilde{y}_i = \tilde{A}_{\circ} \odot \tilde{X}_{i\circ} + \tilde{A}_1 \odot \tilde{X}_{i1} + \cdots + \tilde{A}_p \odot \tilde{X}_{ip}$$

که در آن

$$\tilde{A}_j = (a_j, s_j)_L \quad \tilde{X}_{ij} = (x_{ij}, d_{ij})_L \quad \tilde{y}_i = (y_i, e_i)_L \quad i = 1, 2, \dots, n \quad j = \circ, 1, \dots, p$$

و برای هر دو عدد فازی متقارن $\tilde{A} = (a, \alpha)_L$ و $\tilde{B} = (b, \beta)_L$ عملگر ضرب $\odot$ به صورت زیر خواهد بود:

$$\tilde{A} \odot \tilde{B} = (ab, \alpha\beta)_L$$

در این مسیر آن‌ها با استفاده از رابطه امکان برابری دو عدد فازی [۴۱]، تحت برنامه‌ریزی خطی زیر پارامترهای مدل را بدست آوردند.

$$\begin{aligned} \min Z &= \sum_{i=1}^n \left(\sum_{j=\circ}^p s_j d_{ij} - e_i \right)^2 \\ \text{s. t.} \quad &\sum_{j=\circ}^p a_j x_{ij} + L^{-1}(\alpha) \sum_{j=\circ}^p s_j d_{ij} \geq y_i - L^{-1}(\alpha) e_i \\ &\sum_{j=\circ}^p a_j x_{ij} - L^{-1}(\alpha) \sum_{j=\circ}^p s_j d_{ij} \leq y_i + L^{-1}(\alpha) e_i \\ &\circ \leq \alpha \leq 1 \quad a_j, s_j \in \mathbb{R} \quad j = \circ, 1, \dots, p \quad i = 1, 2, \dots, n, \end{aligned}$$

که در آن $L^{-1}(\alpha)$ معکوس تابع عضویت برای برش دلخواه α است.

مدل نصرآبادی و همکاران (۲۰۰۵)

در سال ۲۰۰۵ نصرآبادی و همکاران [۹۸] با توجه به مشکل اثرگذاری نقاط پرت در برازش مدل رگرسیون فازی با ارائه یک مدل با چند هدف (در واقع توسعه مدل‌های ارائه شده توسط [۱۰۰] و [۱۱۶]) پارامترهای لازم را برآورد کردند و تحت مثال‌هایی عددی اثر بهینگی این روش را نسبت به روش ساکاو و یانو نشان دادند. آن‌ها با فرض مدل به صورت ارائه شده در [۱۱۶] و با فرض LL بودن اعداد فازی موجود در مدل، مرکز و پهنای برآورد شده خروجی را به صورت زیر محاسبه کردند:

$$\hat{y}_i = \hat{a}_0 + \sum_{j=0}^p \hat{a}_j x_{ij}, \quad \hat{e}_i = s_0 + \sum_{j=0}^p (\hat{s}_j |x_{ij}| + d_{ij} |\hat{a}_j|)$$

که با فرض $T = 0, 1, \dots, p$ و $t_1 \cap t_2 = \emptyset$, $t_2 = \{j \in T | a_j \leq 0\}$, $t_1 = \{j \in T | a_j \geq 0\}$ مقدار $E^2 = \sum_{i=1}^n (\varepsilon_{iL}^2 + \varepsilon_{iR}^2)$ هدف با استفاده از تابع هدف (ارائه شده در [۱۰۰]) مدل را به صورت زیر ارائه نمودند:

$$\begin{aligned} \min Z &= \omega \sum_{i=1}^n (\hat{e}_i - e_i)^2 + (1 - \omega) \sum_{i=1}^n (\varepsilon_{iL}^2 + \varepsilon_{iR}^2) \\ \text{s. t. } &\hat{y}_i - y_i \leq \varepsilon_{iL} \\ &y_i - \hat{y}_i \leq \varepsilon_{iR} \\ &e_i \geq 0 \quad \varepsilon_{iL}, \varepsilon_{iR} \geq 0 \quad i = 1, \dots, n \end{aligned}$$

که در آن $0 < w < 1$ توسط فرد محقق تعیین می‌گردد.

مدل مدرس و همکاران (۲۰۰۵)

در مقاله [۹۴] با فرض ورودی غیرفازی، ضرایب و خروجی فازی و توجه به اختلاف بین پهنای برآورد شده و مشاهده شده متغیر پاسخ، (معادل با خطای برآورد در رگرسیون معمولی) تابع هدف را کمترین توانهای دوم تفاضل پهنای کل پاسخ مشاهده شده از پاسخ برآورد شده در نظر گرفتند و شرایط را مانند شرایط موجود در حالت اتصالی [۱۲۹] قرار دادند. البته با این تفاوت که برای تمام اعداد فازی LL ارائه گردیده است. از این‌رو مدل برنامه‌ریزی به صورت

۱۵.۱ ارائه گردید.

$$\begin{aligned}
 \min \quad & Z = \sum_{i=1}^n (s_i |x_{ij}| - e_i)^2 \\
 \text{s. t.} \quad & \sum_{j=0}^p a_j x_{ij} + \left| L^{-1}(\alpha) \right| \sum_{j=0}^p s_j |x_{ij}| \geq y_i - \left| L^{-1}(\alpha) \right| e_i \\
 & \sum_{j=0}^p a_j x_{ij} - \left| L^{-1}(\alpha) \right| \sum_{j=0}^p s_j |x_{ij}| \leq y_i + \left| L^{-1}(\alpha) \right| e_i \\
 & s_j |x_{ij}| \geq 0, \quad j = 0, 1, \dots, p, \quad i = 1, 2, \dots, n.
 \end{aligned} \tag{۱۵.۱}$$

مدل حجتی و همکاران (۲۰۰۵)

در این مقاله [۶۸] دو حالت در نظر گرفته شده است. در حالت اول تنها متغیر ورودی غیر فازی است و در حالت دوم همه عوامل مدل فازی هستند.

در حالت اول آن‌ها با استفاده از روش‌های ارائه شده در [۱۰۰] و [۱۰۳] به ساده کردن فرمول و بهینه کردن آن‌ها پرداختند بدین صورت که با توجه به نتایج هدف ارائه شده در [۱۰۰] مجموع انحرافات بالا و پایین متغیر پاسخ مشاهده شده از برآورد شده (مانند جملات خطا در رگرسیون کلاسیک) را به دو بخش تقسیم کردند. از این‌رو مدل به صورت زیر بازنویسی می‌گردد،

$$\begin{aligned}
 \min \quad & z = \sum_{i=1}^n (d_{iU}^+ + d_{iU}^- + d_{iL}^+ + d_{iL}^-) \\
 \text{s. t.} \quad & \sum_{j=0}^p (a_j + (1 - \alpha)s_j)x_{ij} + d_{iU}^+ - d_{iU}^- = y_i + (1 - \alpha)e_i \\
 & \sum_{j=0}^p (a_j - (1 - \alpha)s_j)x_{ij} + d_{iL}^+ - d_{iL}^- = y_i - (1 - \alpha)e_i \\
 & d_{iU}^+, d_{iU}^-, d_{iL}^+, d_{iL}^-, s_j \geq 0, \quad i = 1, 2, \dots, n, \quad j = 0, 1, \dots, p.
 \end{aligned}$$

که در آن $|d_{iU}^+ - d_{iU}^-|$ فاصله بین نقاط بالایی بازه پیش‌بینی شده با مشاهده شده و همین‌طور $|d_{iL}^+ - d_{iL}^-|$ برای فاصله نقاط پایینی تحت α -برش معلوم است. که در آن d^+ و d^- به عنوان مقدار انحرافات بالایی و پایینی بین متغیرهای پاسخ و برآورد، بر اساس افزودن و یا کاستن مقدار پهنای آن‌هاست. در حالت دوم نیز با استفاده از روش [۱۱۶] و اضافه کردن مقادیر d^+ و d^- جدید، مدل را ارائه کردند.

مدل ترابی و بهبودیان (۲۰۰۷)

در مقاله [۱۳۵] برای مواجهه با داده‌های دور افتاده ترابی و بهبودیان با استفاده از روش حداقل قدر مطلق انحرافات فازی در مدلی که دارای جزء خطای فازی بود و براساس اصل تجزیه تحت

چهار نوع برنامه‌ریزی خطی با توابع هدف متفاوت با ورودی و خروجی فازی و ضرایب فازی پارامترهای مجهول را تحت α - برش‌های مختلف بدست آوردند و نتایج را برای انواع مدل‌های برنامه‌ریزی خطی خود تحت مثال‌های مختلف نشان دادند.

مدل حسن‌پور و همکاران (۲۰۰۹)

پس از کیم-بیشو در [۸۲]، حسن‌پور و همکاران در [۶۰] با ارائه مثالی نشان دادند که روش ارائه شده [۸۲] می‌تواند به‌گونه‌ای ترتیب چپ و راست بودن مقادیر پهناها را نسبت به مرکز رعایت نکند. از این‌رو برای بهبود این روش ابتدا از همان روش ارائه شده [۸۲] مراکز ضرایب را برآورد کرده سپس مدل برنامه‌ریزی ۱۶.۱ را ارائه نمودند. همچنین آن‌ها طی مثالی نشان دادند در مواردی که روش [۸۲] بدون ایراد است مدل آن‌ها جواب‌های مشابه ارائه می‌نماید.

$$\begin{aligned} \min \quad & Z = \sum_{i=1}^n \left(\sum_{j=0}^p l_j x_{ij} - \left(y_i - |L^{-1}(\alpha)| e_i \right) \right)^2 \\ \text{s. t.} \quad & l_j \leq a_j \quad j = 0, 1, \dots, p \\ \min \quad & Z = \sum_{i=1}^n \left(\sum_{j=0}^p r_j x_{ij} - \left(y_i - |L^{-1}(\alpha)| e_i \right) \right)^2 \\ \text{s. t.} \quad & r_j \geq a_j \quad j = 0, 1, \dots, p \end{aligned} \quad (16.1)$$

که در آن، ورودی‌ها غیر فازی و خروجی و ضرایب فازی LR فرض شده است و داریم،

$$l_j = a_j - s_j, \quad r_j = a_j + s'_j$$

که s_j و s'_j به ترتیب پهناهای چپ و راست عدد فازی LR هستند.

مدل حسن‌پور و همکاران (۲۰۱۰)

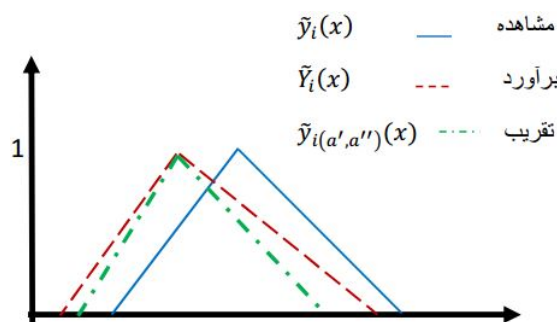
حسن‌پور و همکاران در [۶۱] مدلی با ورودی و خروجی فازی LR و ضرایب غیرفازی را در نظر گرفتند به‌طوری‌که به ازای هر i و j ، $\tilde{y}_i = (y_i, l_i, r_i)$ و $\tilde{x}_{ij} = (x_{ij}, l_{ij}, r_{ij})$ که در آن l_{ij} ، r_i و r_{ij} و l_i به ترتیب مقادیر بالا و پایین اعداد فازی هستند. آن‌ها با استناد به اینکه هر عدد حقیقی a را می‌توان به صورت تفاضل دو عدد حقیقی مثبت به صورت $a = a' - a''$ نوشت، مدل را بصورت زیر ارائه کردند،

$$\tilde{Y}_i \subseteq \hat{\tilde{Y}}_i = \left(a_0 + \sum_{j=1}^p (a'_j - a''_j) x_{ij}, \sum_{j=1}^p (a'_j l_{ij} + a''_j r_{ij}), \sum_{j=1}^p (a'_j r_{ij} - a''_j l_{ij}) \right)$$

در این صورت مدل برنامه‌ریزی خطی به صورت زیر ارائه گردید.

$$\begin{aligned}
 \min \quad & Z = \sum_{i=1}^n (n_{ic} + p_{ic} + n_{il} + p_{il} + n_{ir} + p_{ir}) \\
 \text{s. t.} \quad & a_0 + \sum_{j=1}^p (a'_j - a''_j) x_{ij} + n_{ic} - p_{ic} = y_i \\
 & \sum_{j=1}^p (a'_j l_{ij} + a''_j r_{ij}) + n_{il} - p_{il} = l_i \\
 & \sum_{j=1}^p (a'_j r_{ij} - a''_j l_{ij}) + n_{ir} - p_{ir} = r_i \\
 & a_0 \in \mathbb{R}, \quad a'_j, a''_j \geq 0, \quad n_{ik}, p_{ik} \geq 0, \quad n_{ik} p_{ik} = 0 \\
 & i = 1, 2, \dots, n, \quad j = 1, 2, \dots, p, \quad k = l, c, r
 \end{aligned} \tag{۱۷.۱}$$

که در آن، با فرض $\tilde{Y}_i(a', a'') = (y_i(a', a''), l_i(a', a''), r_i(a', a''))$ داریم،
 n_{ic} و p_{ic} به ترتیب عبارتند از انحرافات مثبت و منفی بین مرکز پاسخ مشاهده شده و برآورد
 شده l_i و همین‌طور p_{il} و n_{il} نیز به ترتیب انحرافات مثبت و منفی بین پهنای
 چپ (راست) آن‌ها هستند.



و همین‌طور کار برای عدد فازی دوزنقه‌ای هم انجام گردیده است. همچنین آن‌ها برای
 مقایسه دو عدد فازی LR متر ۱۸.۱ را ارائه کردند و به کار بردند:

$$d(\tilde{A}_1, \tilde{A}_2) = |a_1 - a_2| + |r_1 - r_2| + |l_1 - l_2| \tag{۱۸.۱}$$

که در آن $\tilde{A}_1 = (a_1, l_1, r_1)_{LR}$ و $\tilde{A}_2 = (a_2, l_2, r_2)_{LR}$ اعداد فازی LR می‌باشند.

مدل حسن‌پور و همکاران (۲۰۱۱)

در سال ۲۰۱۱، حسن‌پور و همکاران [۶۲] مدلی را بررسی کردند که تمام عوامل آن اعداد فازی
 بودند. آن‌ها با توجه به اینکه ضرب دو عدد فازی LR به‌طور تقریبی یک عدد فازی LR می‌شود
 و براساس روش ارائه شده در [۶۱] مدل را برآورد کرده و تحت متر ارائه شده خود در [۶۱] به
 ارزیابی مدل نسبت به دیگر روش‌ها پرداختند.

مدل چن و همکاران (۲۰۱۳)

در این مدل [۳۲]، که تعمیم دیدگاه کائو و چیو [۷۸] می‌باشد، ورودی و خروجی فازی و ضرایب غیرفازی است. در مرحله اول بر اساس آلفا-برش‌ها، حدود بالا و پایین متغیر پاسخ برآورد شده را بر اساس ضرایب غیر فازی و حدود بالا و پایین متغیرهای توضیحی متناظر با آن α - برش تعیین می‌شود. سپس در مرحله دوم بر اساس می‌نیمم کردن یک متر و از طریق مدل برنامه‌ریزی خطی زیر این ضرایب غیر فازی را برآورد کردند و با مثال‌هایی برتری روش خود را نشان دادند.

$$\begin{aligned} \min \quad & \sum_{i=1}^n \frac{1}{2m} \sum_{k=1}^m \left(\left| (\hat{y}'_i)_{\alpha_k}^L - (y_i)_{\alpha_k}^L \right| + \left| (\hat{y}'_i)_{\alpha_k}^U - (y_i)_{\alpha_k}^U \right| \right) \\ \text{s.t.} \quad & (\hat{y}'_i)_{\alpha_k}^L = b_0 + b_1(X_{i1})_{\alpha_k}^L + b_2(X_{i2})_{\alpha_k}^L + \dots + b_p(X_{ip})_{\alpha_k}^L + (\delta)_{\alpha_k}^L \\ & (\hat{y}'_i)_{\alpha_k}^U = b_0 + b_1(X_{i1})_{\alpha_k}^U + b_2(X_{i2})_{\alpha_k}^U + \dots + b_p(X_{ip})_{\alpha_k}^U + (\delta)_{\alpha_k}^U \\ & i = 1, \dots, n \quad k = 1, \dots, m \end{aligned}$$

که در آن $\hat{y}'_i$ پاسخ برآورد شده در مرحله اول است.

مدل ربیعی و همکاران (۲۰۱۵)

در این مدل [۱۰۹] ربیعی و همکاران با استفاده از اعداد فازی بازه‌ای-مقدار مثلثی و تعریف یک فاصله روی فضای اعداد بازه‌ای-مقدار مثلثی، به برآورد پارامترهای مدل خود پرداختند.

مدل حسین زاده و همکاران (۲۰۱۶)

این مدل [۷۲] بر اساس روش برنامه ریزی خطی ارائه شده در حسن پور و همکاران [۶۱، ۶۲] عمل می‌کند با این تفاوت که متغیر پاسخ مورد بحث یک عدد فازی نوع ۲ است. البته این مقاله دقیقاً مانند [۱۰۸] با این تفاوت که اعداد فازی نوع ۲ و از روش برنامه ریزی هدف وزن‌دار استفاده شده است.

مدل کرنی و هالدیک (۲۰۱۸)

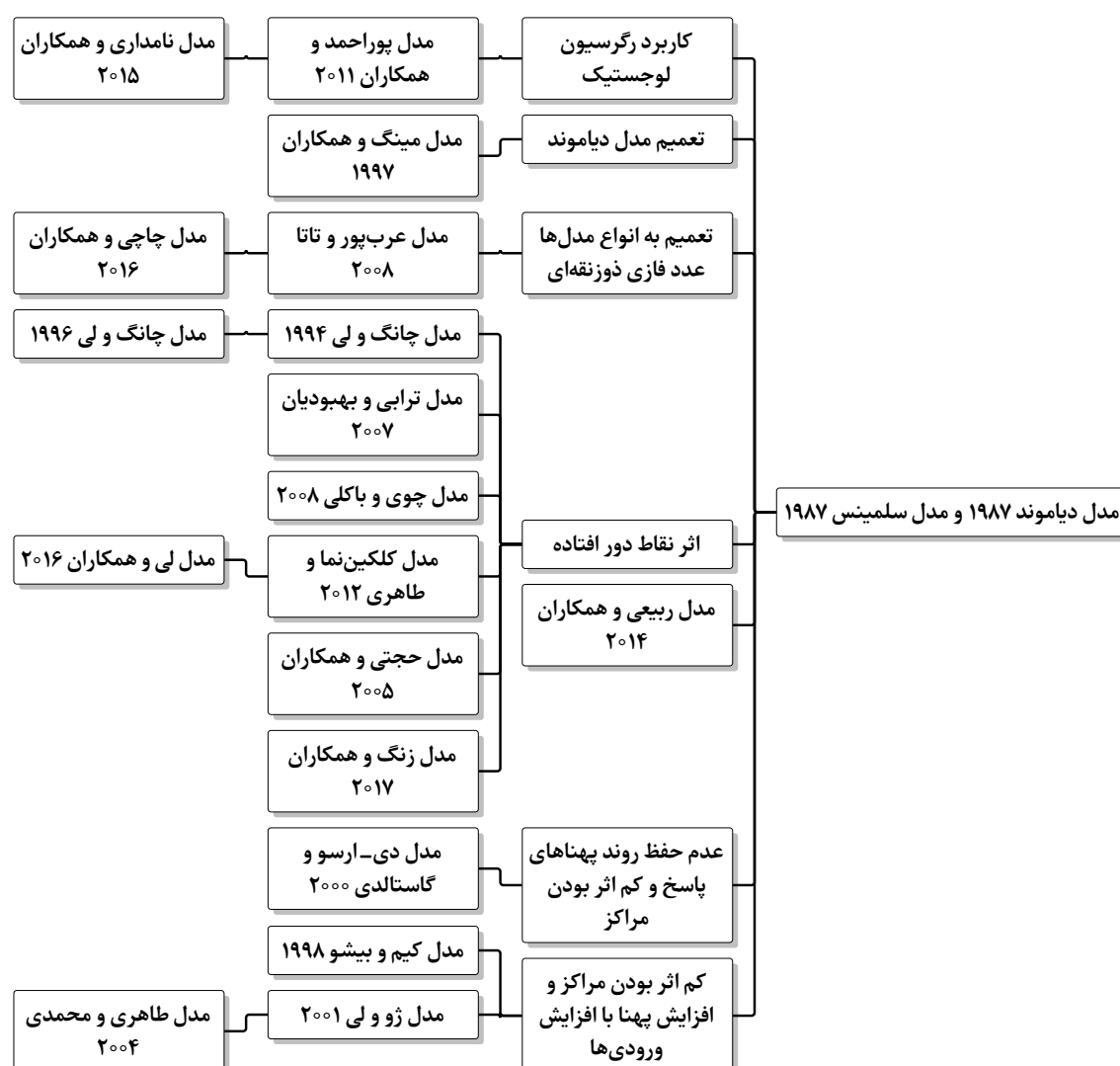
در این مدل [۲۴] دیدگاه امکانی برای مدل بندی داده‌هایی که ضرایب آن‌ها بر اساس یک سری اطلاعات پیشین دارای محدودیت مقدار می‌باشند بکار گرفته شد. در واقع با استفاده از تعریف عملگر بازه‌ای، برای خروجی برآورد شده بازه‌ای تولید کردند و با کمینه کردن مجموع ضرایب محدود شده طی یک مدل برنامه ریزی پارامترهای مدل را برآورد کردند. مدل پیشنهادی بسیار گسترده است و حتی برای ورودی‌های فازی بازه‌ای-مقدار هم ارائه شده است اما آن‌ها

برای نمایش عملکرد روش خود از اعداد فازی مثلثی و نرمال برای مقادیر ورودی و خروجی استفاده کردند.

۴.۱ مدل‌های کمترین فاصله

این روش‌ها بر اساس کمینه کردن فاصله بین تابع عضویت پاسخ‌های مشاهده شده و برآورد شده عمل می‌کند. در واقع با استفاده از یک متر و کمینه کردن فاصله دو عدد فازی، پارامترهای مدل بدست می‌آید. این مسئله در روش تاناکا و همکاران وجود ندارد. نمودار درختی مربوط به مدل‌های کمترین فاصله به صورت زیر خواهد بود.

درخت رشد رگرسیون فازی در رویکرد کمترین فاصله



شکل ۵.۱: درخت رشد رگرسیون فازی در رویکرد کمترین فاصله

مدل سلمینس (۱۹۸۷)

پس از مدل [۱۵]، سلمینس [۲۳] با ارائه یک اندازه به صورت

$$\gamma(\tilde{A}, \tilde{B}) = \max_x \min \{ \tilde{A}(x), \tilde{B}(x) \}$$

بنام اندازه سازگاری، با فرض اینکه $\tilde{A}$ و $\tilde{B}$ دو عدد فازی مثلثی باشند به ارائه یک مدل رگرسیونی پرداخت. در این دیدگاه با توجه به اینکه $0 \leq \gamma \leq 1$ است هرچه γ به صفر نزدیک شود همپوشانی دو عدد فازی کمتر خواهد شد و هرچه به یک نزدیک شود همپوشانی افزایش می‌یابد. در واقع هدف، پیدا کردن مدلی است که سازگاری کلی بین داده‌ها و مدل برازش شده ماکسیمم شود پس تابع هدف عبارت است از،

$$\begin{aligned} \hat{\tilde{Y}} &= \tilde{A}_0 + \tilde{A}_1 x \\ &= a_0 + a_1 x \pm \sqrt{s_0^2 + 2s_0 s_1 x + s_1^2 x^2} \end{aligned}$$

که در آن $\tilde{A}_1 = (a_1, s_1)_T$ و $\tilde{A}_0 = (a_0, s_0)_T$ هستند و مراکز اعداد فازی برآورد شده $\tilde{A}_1$ و $\tilde{A}_0$ براساس رگرسیون کمترین توان دوم موزون بدست می‌آید و s_0 با عنوان هماهنگی فازی بین ضرایب فرض شده است که مفهوم مشابهی با کوواریانس بین دو پارامتر را در رگرسیون معمولی دارد.

مدل دیاموند (۱۹۸۷)

این مدل در سال ۱۹۸۷ توسط دیاموند [۴۰] ارائه گردید که با ایده گرفتن از روش کمترین توانهای دوم در رگرسیون معمولی عمل می‌کند. در این مسیر یک بار ورودی و خروجی مدل عدد فازی مثلثی و ضرایب آن دقیق، و بار دیگر ضرایب و خروجی مدل عدد فازی مثلثی و ورودی‌های آن دقیق بودند. او براساس فاصله بین دو نقطه، فاصله بین عدد فازی پاسخ و برآورد را به صورت زیر در نظر گرفت،

$$D^2(\tilde{Y}_i, \hat{\tilde{Y}}_i) = [(y_i - \hat{y}_i)^2 + (y_i - \hat{y}_i - (e_i - \hat{e}_i))^2 + (y_i - \hat{y}_i + (e_i - \hat{e}_i))^2] \quad (19.1)$$

دارا بودن کارائی بالاتر در پیش‌بینی از مزیت‌های این روش است. اما ضعف این روش عدم وجود پوشش برای تمامی اعداد فازی است. البته پیچیدگی روش‌های حداقل مربعات نسبت به روش‌های امکانی نیز با افزایش متغیرهای ورودی و تنوع در فازی بودن عوامل مدل نیز قابل اشاره است.

مدل چانگ و لی (۱۹۹۴)

چانگ و لی [۳۰] با استفاده از روش رتبه‌بندی اعداد فازی، هم به می‌نیمم کردن توان دوم اعداد فازی مشاهده شده و برآورد شده و هم می‌نیمم کردن قدرمطلق تفاضل رتبه این اعداد پرداختند. در مدل بکار رفته داده‌های ورودی حقیقی است ولی ضرایب و خروجی مدل عدد فازی LR فرض شده است. این روش در دو حالت مثلثی و درجه دوم ارائه گردید.

همچنین در سال ۱۹۹۶ [۱۱۲] رگرسیون کمترین توانهای دوم فازی وزن دار شده را ارائه کردند. آن‌ها به این مطلب اشاره داشتند که در روش‌های پیشنهاد شده قبلی از دیدگاه کمترین توانهای دوم، بحث اثرات متقابل مراکز و پهناها در نظر گرفته نشده و فرایند تحلیل آن‌ها جداگانه انجام شده است و مزیت روش خود را در سه موضوع به شرح زیر بیان نمودند:

الف) بهبود برآورد پارامترهای فازی با اثر متقابل یا بدون آن،

ب) اطمینان محقق در مورد تجمیع داده‌ها و مدل طبقه‌بندی شده در قالب یک فرایند یکپارچه،

ج) کاهش قابل ملاحظه اثرات نقاط دور افتاده موجود در مشاهدات.

مدل پیشنهادی با ورودی و خروجی دقیق بحث شده که تنها ضرایب مدل عدد فازی LR فرض شده‌اند به صورت زیر است،

$$\begin{aligned}\min Z_M &= (\mathbf{Y} - \mathbf{X}\mathbf{a})^\top \mathbf{F}^M (\mathbf{Y} - \mathbf{X}\mathbf{a}) \\ \min Z_L &= ((\mathbf{X}\mathbf{a} - \mathbf{Y}) - \mathbf{X}\mathbf{S}^{L,\alpha})^\top \mathbf{F}^L ((\mathbf{X}\mathbf{a} - \mathbf{Y}) - \mathbf{X}\mathbf{S}^{L,\alpha}) \\ \min Z_R &= ((\mathbf{Y} - \mathbf{X}\mathbf{a}) - \mathbf{X}\mathbf{S}^{R,\alpha})^\top \mathbf{F}^R ((\mathbf{Y} - \mathbf{X}\mathbf{a}) - \mathbf{X}\mathbf{S}^{R,\alpha})\end{aligned}\quad (20.1)$$

که در آن $\mathbf{X}$ و $\mathbf{Y}$ ماتریس مشاهدات تحت نمونه n تایی است. $\mathbf{F}^M$ ، $\mathbf{F}^L$ و $\mathbf{F}^R$ ماتریس‌های قطری $p \times p$ با عناصر وزنی مربوط به هر نمونه هستند و همچنین داریم:

$$\begin{aligned}\mathbf{a} &= (\mathbf{X}^\top \mathbf{F}^M \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{F}^M \mathbf{Y} \\ \mathbf{S}^{L,\alpha} &= (\mathbf{X}^\top \mathbf{F}^L \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{F}^L (\mathbf{X}\mathbf{a} - \mathbf{Y}) \\ \mathbf{S}^{R,\alpha} &= (\mathbf{X}^\top \mathbf{F}^R \mathbf{X})^{-1} \mathbf{X}^\top \mathbf{F}^R (\mathbf{Y} - \mathbf{X}\mathbf{a})\end{aligned}\quad (21.1)$$

علاوه بر آن مزیت دیگر این روش با استفاده از یک معیار نیکویی برازش، بحث انتخاب متغیر است. بدین معنی که می‌توان روش‌هایی مانند الگوی جلورونده^۱، عقب‌رونده^۲ و قدم به قدم^۳ را در آن بررسی کرد.

^۱Forward

^۲Backward

^۳Stepwise

مدل مینگ و همکاران (۱۹۹۷)

دیاموند در [۴۰] حالت خاصی از اعداد فازی را مورد بررسی قرار داد و این باعث گردید تا مینگ و همکاران [۹۲] تعمیم آن را ارائه نمایند. در واقع روش دیاموند تنها برای اعداد فازی مثلثی پیشنهاد شده بود، اما در این مقاله برای تمام اعداد فازی روش کمترین توان دوم فازی ارائه شده است.

مدل کیم و بیشو (۱۹۹۸)

در سال ۱۹۹۸ با استناد به اینکه در روش تاناکا و همکاران [۱۲۸] اولاً برای پوشش α -برش پاسخ مشاهده شده توسط α -برش پاسخ برآورد شده لازم است پهنای برآورد شده زیاد باشد، دوماً نقاط با درجات عضویت بالا کاربرد مفیدی ندارند، سوماً با افزایش متغیر مستقل (ورودی)، پهنای برآورد شده برای پاسخ افزایش می‌یابد، کیم و بیشو [۸۲] با فرض مدل به صورت ورودی حقیقی، ضرایب و خروجی فازی (عدد فازی LR) در سه مدل رگرسیونی مجزا برای مرکز، پاسخ، مقدار پاسخ چپ و مقدار پاسخ راست، تحت روش رگرسیونی کمترین توان‌های دوم کلاسیک، برآوردهای پارامترها را به‌دست آوردند. برای حالت عدد فازی مثلثی متقارن معادلات پایه به صورت ۲۲.۱ خواهد بود

$$\begin{aligned} y_i - (1 - \alpha)e_i &= \sum_{j=0}^p s_j x_{ij} \\ y_i &= \sum_{j=0}^p a_j x_{ij} \\ y_i + (1 - \alpha)e_i &= \sum_{j=0}^p s_j x_{ij} \end{aligned} \quad (22.1)$$

البته آن‌ها در مقاله خود α را برابر صفر فرض کردند.

مدل دورسو و گاستالدی (۲۰۰۰)

در سال ۲۰۰۰ یک دیدگاه جدیدی در رگرسیون حداقل مربعات فازی در [۴۳] با عنوان مدل رگرسیون فازی انطباقی^۴ خطی دوگانه ارائه گردید که بر اساس دو مدل خطی کار کرد. یک مدل رگرسیونی مراکز و یک مدل رگرسیونی پهنایها.

در این روش با این دید که یک رابطه‌ای بین تعیین مقدار پهنایها از روی مراکز در واقعیت وجود دارد (مثلاً درصدی از مرکز را به عنوان پهنایها در نظر می‌گیرند) عمل کرده و برآورد مرکز

^۴ adaptive fuzzy regression

و پهنای پاسخ مشاهدات را به صورت زیر فرض کردند،

$$\mathbf{y} = \hat{\mathbf{y}} + \varepsilon_y, \quad \hat{\mathbf{y}} = \mathbf{X}\mathbf{a} \quad (23.1)$$

$$\mathbf{e} = \hat{\mathbf{e}} + \varepsilon_e, \quad \hat{\mathbf{e}} = \mathbf{X}\mathbf{a}b + \mathbf{1}d \quad (24.1)$$

که در آن $\mathbf{X}$ ماتریسی با ابعاد $n \times (p+1)$ است (ماتریس مشاهدات)، $\mathbf{a}$ بردار با $(p+1)$ ستون شامل پارامترهای رگرسیون در مدل مربوط به مراکز، $\mathbf{y}$ و $\hat{\mathbf{y}}$ به ترتیب مراکز و برآورد آن در پاسخ، $\mathbf{e}$ و $\hat{\mathbf{e}}$ پهنای متناظر با آن‌هاست. b و d اعداد حقیقی‌اند که پارامترهای مدل رگرسیونی مربوط به پهنای هستند و $\mathbf{1}$ یک بردار $n \times 1$ با مقادیر ۱ است. تابع هدف نیز براساس فاصله اقلیدسی $\hat{\mathbf{y}}$ و $\tilde{\mathbf{y}}$ به ترتیب به عنوان بردار پاسخ برآورد شده و مشاهده شده هستند. است که می‌بایست می‌نیمم گردد. یعنی داریم،

$$Z = \sum_{i=1}^n D_i^2 = (\mathbf{y} - \hat{\mathbf{y}})^\top (\mathbf{y} - \hat{\mathbf{y}}) + (\mathbf{e} - \hat{\mathbf{e}})^\top (\mathbf{e} - \hat{\mathbf{e}}) \quad (25.1)$$

که برای می‌نیمم کردن Z با مشتق‌گیری نسبت به پارامتر فوق‌الذکر، برآورد آن‌ها عبارتند از،

$$\begin{aligned} \mathbf{a} &= \frac{1}{(1+b^2)} \left((\mathbf{X}^\top \mathbf{X})^{-1} \mathbf{X}^\top (\mathbf{y} + \mathbf{e}b - \mathbf{1}bd) \right) \\ b &= (\mathbf{a}^\top \mathbf{X}^\top \mathbf{X} \mathbf{a})^{-1} (\mathbf{e}^\top \mathbf{X} \mathbf{a} - \mathbf{a}^\top \mathbf{X}^\top \mathbf{1}d) \\ d &= \frac{1}{n} \left(\mathbf{e}^\top \mathbf{1} - \mathbf{a}^\top \mathbf{X}^\top \mathbf{1}b \right) \end{aligned} \quad (26.1)$$

در این مدل خواص زیر برقرار می‌گردد،

الف) مجموع n واحد خطا بین مرکز پاسخ مشاهده شده و برآورد شده صفر است.

$$\mathbf{1}^\top (\mathbf{y} - \hat{\mathbf{y}}) = 0 \quad (27.1)$$

ب) مجموع n واحد خطا بین پهنای پاسخ مشاهده شده و برآورد شده صفر است.

$$\mathbf{1}^\top (\mathbf{e} - \hat{\mathbf{e}}) = 0 \quad (28.1)$$

ج) عبارت‌های زیر برقرار است،

$$(\mathbf{e} - \hat{\mathbf{e}})^\top \hat{\mathbf{y}} = 0, \quad (\mathbf{y} - \hat{\mathbf{y}})^\top \hat{\mathbf{e}} = 0 \quad (29.1)$$

همچنین در سال (۲۰۰۲) با وزن‌دار کردن مراکز و پهنای در متر مبنا، پاسخ‌های برآورد شده را تحت مدل‌های چند جمله‌ای در [۴۴] ارائه دادند. همچنین صورت بسته ضرایب تولید کننده مدل را برای خروجی با تابع عضویت ذوزنقه‌ای تولید نمودند.

مدل ژو و لی (۲۰۰۱)

در رگرسیون کلاسیک روش کمترین توانهای دوم عمومی‌ترین روش بدست آوردن ضرایب مجهول است. لذا ژو و لی [۱۴۲] با استفاده از این دیدگاه با فرض یک مدل با ورودی‌های غیرفازی مثبت و خروجی و ضرایب فازی نرمال برای مرکز و پهنای ضرایب مجهول به صورت ۳۰.۱ برآورد را بدست آوردند.

$$\hat{\mathbf{a}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{y} \quad \hat{\mathbf{s}} = (\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{e} \quad (30.1)$$

که در آن $\mathbf{X}$ ماتریس مقادیر ورودی و $\mathbf{y}$ و $\mathbf{e}$ به ترتیب ماتریس مراکز و پهنای متغیر پاسخ است. در این مقاله فرض اصلی بر آن است که مقادیر $\hat{\mathbf{s}}$ همگی مثبت باشد بدین معنی که ماتریس $(\mathbf{X}^T \mathbf{X})^{-1} \mathbf{X}^T \mathbf{e}$ معین مثبت باشد.

مدل دورسو (۲۰۰۳)

دورسو در تعمیم کارهای خود، در [۴۲] فرض فازی بودن داده‌های ورودی را برای توابع عضویت مثلثی و دوزنقه‌ای در نظر گرفت. علاوه بر آن برای مشکل منفی شدن پهنای برآورد شده پیشنهاداتی را نیز ارائه کرد که از جمله آن که در صورت منفی برآورد شدن پهنای، مقدار صفر برای آن در نظر گرفته شود.

مدل محمدی و طاهری (۲۰۰۴)

در مقاله [۱۴۲] برای حالاتی که برآورد پهنای ضرایب منفی شود بحث نگردیده است. لذا محمدی و طاهری [۹۳] برای رفع این مسئله این ایده را ارائه کردند که در صورت منفی شدن هر کدام از برآوردهای پهنای، آن را صفر فرض کرده و از رابطه (۳۰.۱) دوباره برای بقیه پهنایها برآورد جدید ارائه کردند. این کار تا زمانی که برآوردهای ارائه شده همگی مثبت شوند ادامه می‌یابد.

مدل کی و همکاران (۲۰۰۶)

با استناد به روش دورسو و استفاده از متر [۱۴۴]، برای اعداد فازی LR در [۳۶] مدل خود را ارائه کردند. آن‌ها با اشاره به مسئله همخطی در ماتریس طرح با بررسی ضریب تعیین تعدیل شده در مثال‌های خود شیوه انتخاب متغیر را نیز اجرا نمودند.

مدل چوی و باکلی (۲۰۰۸)

در بررسی مدل‌های رگرسیون فازی، یکی از مشکلات ایجاد شده مواجه با داده‌های دور افتاده است که می‌تواند خطای برآورد را افزایش دهد. از این‌رو چوی و باکلی [۳۴] با استفاده از

روش‌های بکار رفته توسط [۸۲] و [۷۷] و جایگذاری قدرمطلق بجای توان‌های دوم در آن‌ها ایده خود را ارائه نموده تحت مثال‌هایی روش‌های فوق را با روش خود مقایسه کرده و بهبود خطای کل را نشان دادند.

مدل عرب‌پور و تاتا (۲۰۰۸)

دیاموند در [۴۰] از طریق رویکرد کمترین توان‌های دوم به برآورد پارامترهای مدل با فرض مثلی بودن اعداد فازی بکار رفته پرداخت. در این مسیر عرب‌پور و تاتا [۱۳] با ارائه یک متر، از طریق روش کمترین توان‌های دوم برای سه وضعیت:

(الف) ورودی و خروجی فازی و ضرایب غیرفازی

(ب) ورودی غیرفازی و خروجی و ضرایب فازی

(ج) ورودی و خروجی و ضرایب فازی

که همگی با در نظر گرفتن وجود خطای فازی در مدل بوده است. برای اعداد فازی مثلی و دوزنقه‌ای و همچنین برای حالت چند متغیره با اعداد فازی دوزنقه‌ای به برآورد پارامترها پرداختند و تحت مثال‌هایی خطای کل مدل را با روش‌های مشابه مقایسه نمودند.

مدل طاهری و کلین‌نما (۲۰۱۲)

بدنبال مشکل تاثیرپذیری از داده‌های دور افتاده، طاهری و کلین‌نما در [۸۱] با ارائه یک متر قدرمطلق به صورت $D_{LR}(\tilde{M}, \tilde{Q})$ تحت عملگرهای $\oplus_w$ و $\odot_w$ تحت مدل برآوردشده

$$\hat{\tilde{y}} = \tilde{M} \oplus_w (\tilde{Q} \odot_w \tilde{X})$$

با فرض

$$\tilde{M} \oplus_w \tilde{Q} = (m + q, \max\{\alpha_M q, \alpha_Q m\}, \max\{\beta_M q, \beta_Q m\})_{LR}$$

$$\tilde{M} \odot_w \tilde{Q} = (mq, \max\{\alpha_M q, \alpha_Q m\}, \max\{\beta_M q, \beta_Q m\})_{LR}$$

(به جهت اختصار فرمول $\odot_w$ را فقط برای $m > 0$ و $q > 0$ نوشته‌ایم)

$$D_{LR}(\tilde{M}, \tilde{Q}) = \frac{1}{3} \{ |m - q| + |(m - l\alpha_M) - (Q - l\alpha_Q)| + |(M - r\beta_M) - (Q - r\beta_Q)| \}$$

که در آن $l = \int_0^1 L^{-1}(w)dw$ و $r = \int_0^1 R^{-1}(w)dw$ است، ضرایب مجهول را یکبار با استفاده از خطای فازی و ضرایب غیرفازی و یکبار با استفاده از ضرایب فازی بدست آورده تحت مثال‌هایی بهینگی آن‌ها را نشان دادند.

مدل نامداری و همکاران (۲۰۱۵)

این مدل از رویکرد کمترین قدر مطلق انحراف در حوزه مدل بندی رگرسیون لجستیک فازی استفاده شده است و با استفاده از ورودی غیرفازی ضرایب فازی مثلثی ارائه شده است. در این مقاله یک شاخص نیکویی برازش جدید به نام اندازه عملکرد بر اساس فاصله فازی و متناظر با آن یک شاخص حساسیت ارائه شده است. نتایج حاصل از این روش با نتایج روش کمترین توان دوم خطا در مثال‌هایی واقعی مقایسه شده است. لازم به ذکر است که در سال ۲۰۱۱ نیز این کاربرد در [۱۰۵] بر اساس روش کمترین توان دوم خطا توسط پورمند و همکاران در مطالعات بالینی بکار رفته است.

مدل لی و همکاران (۲۰۱۶)

در این روش [۸۶] مدل رگرسیونی برای اعداد فازی ذوزنقه‌ای ارائه شده است. در این مدل بر اساس اصل توسیع عملگرهای جمع، تفریق و ضرب در اعداد فازی ذوزنقه‌ای تعریف گردید. آن‌ها روش خود را در سه حالت مختلف که در آن مدل با ضرایب غیر فازی، با ورودی غیر فازی و تماماً فازی است ارائه نمودند. این مدل با حداقل کردن مجموع انحرافات فاصله مراکز و پهنای از یکدیگر به دنبال برآورد ضرایب مدل بودند و برتری روش خود را با مثال‌هایی به نمایش گذاشتند. نکته قابل توجه در این روش این است مدل ارائه شده آن‌ها تغییر علامت عوامل را نیز لحاظ نموده‌اند.

مدل زنگ و همکاران (۲۰۱۷)

در این مدل [۱۵۰] اعداد فازی مثلثی برای خروجی و ضرایب در نظر گرفته شده است اما متغیرهای توضیحی غیرفازی هستند. این مدل با استفاده از روش حداقل انحرافات به برآورد پارامترهای مدل پرداخته است. نکته متمایز کننده در این مقاله از روش‌های قبلی مبتنی بر معیار کمترین فاصله استفاده از یک معیار شباهت بود که بر مبنای متر کمترین انحراف است. بدین صورت که در شرایط برآورد، علاوه بر مثبت فرض کردن پهنای ضرایب، مثبت بودن پهنای متغیر پاسخ برآورد شده را نیز جزء شرایط مدل قرار داده است. این مسئله باعث می‌شود که روش برآورد یابی تطابق بیشتری با پاسخ‌های مشاهده شده داشته باشد. همچنین آن‌ها استوار بودن مدل پیشنهادی خود را در برخورد با داده‌های دورافتاده در حالت‌های مختلف با ارائه مثالی عددی نشان دادند. بنابراین با فرض پاسخ مشاهده شده i ام به صورت $\tilde{y}_i = (y_i, l_i, r_i)_T$ با ماتریس متغیرهای ورودی X و ضریب $\tilde{\beta}_j = (c_j, a_j, b_j)_T$ $j = 0, 1, \dots, p$ به

برآورد ضرایب از طریق رابطه ۳۱.۱ پرداختند.

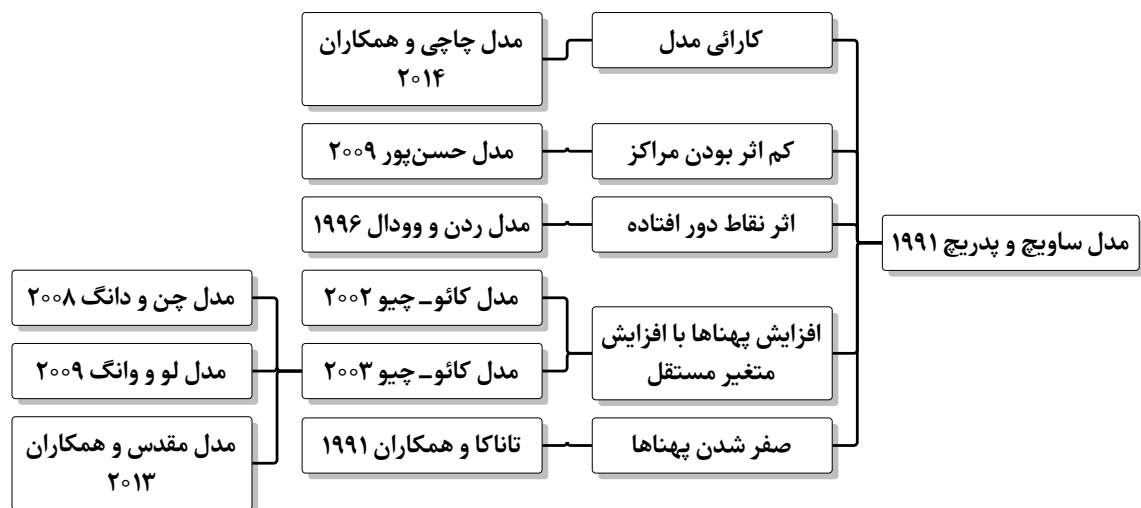
$$\min d = \sum_{i=1}^n |\tilde{y}_i - \hat{y}_i| = \sum_{i=1}^n \left(\left| y_i - \sum_{j=0}^p c_j x_{ij} \right| + \left| l_i - \sum_{j=0}^p a_j x_{ij} \right| + \left| y_i - \sum_{j=0}^p b_j x_{ij} \right| \right) \quad (31.1)$$

$$\text{s.t.} \quad a_j \geq 0, b_j \geq 0, \sum_{j=0}^p a_j x_{ij} \geq 0, \sum_{j=0}^p b_j x_{ij} \geq 0$$

۵.۱ مدل‌های ترکیبی

در این روش‌ها می‌توان گفت که حل مسئله در یک فرآیند دو یا چند مرحله‌ای اجرا می‌گردد. مرحله‌ای بر اساس روش‌های امکانی و مرحله دیگر بر اساس روش کمترین توانهای دوم انجام می‌شود. در اینجا، ابتدا درخت رشد مدل‌های ترکیبی را قبل از توضیح و بیان مدل‌های مربوطه نشان می‌دهیم.

درخت رشد رگرسیون فازی در رویکرد ترکیبی



شکل ۶.۱: درخت رشد رگرسیون فازی در رویکرد ترکیبی

مدل ساویچ و پدريچ (۱۹۹۱)

این روش در سال ۱۹۹۱ ارائه گردید [۱۲۰] و دلیل ارائه آن این بود که در روش مدل‌بندی تاناکا [۱۲۸] با اینکه مراکز دارای بالاترین درجه عضویت هستند ولی اثر آن‌ها در تعیین ضرایب مدل کم است یعنی با وجود اینکه $\tilde{Y}_i \subseteq \tilde{Y}$ است ولی ممکن است مراکز آنها فاصله نسبتاً زیادی

باهم داشته باشد. از این رو در مرحله اول، مراکز ضرایب را از طریق روش کمترین توان دوم، رگرسیون کلاسیک، محاسبه نموده و در مدل برنامه‌ریزی خطی تاناکا [۱۲۸] با جایگذاری آن‌ها در شرایط مدل، پهنای ضرایب را به‌دست آوردند.

مدل تاناکا و همکاران (۱۹۹۱)

پس از روش ارائه شده در [۱۲۹]، تاناکا [۱۳۰] در سال ۱۹۹۱ برای رفع مشکل مدل خود، تابع هدف درجه دومی را برای یک مدل چند متغیره ارائه و آن را در سه حالت مسئله‌ای می‌نیمم، ماکزیمم و اتصال (قطع کردن) بسط داد. بدین شکل مدل‌های ارائه شده براساس یک برنامه‌ریزی غیرخطی کار خواهند کرد.

مدل ردن و وودال (۱۹۹۶)

در سال ۱۹۹۶، ردن و وودال [۱۱۲] با بررسی مقاله ساکاو و یانو [۱۱۶] به این مطلب اشاره کردند که این روش اساساً یک روش آزمون و خطا است. در واقع مثالی ارائه کردند که در آن روش ذکر شده دارای جواب یکتایی نیست بواسطه آنکه قیدهای مسئله قیدهای همپوشانی خواهند داشت. از این رو پیشنهاد کردند که مراکز ضرایب را از طریق رگرسیون کمترین توان‌های دوم بدست آورده و آن‌ها را در فرمول‌های ۷.۱ قرار دهند.

مدل کائو-چیو (۲۰۰۲)

در روش‌ها و مدل‌های قبل عموماً ضرایب رگرسیون اعداد فازی فرض شده‌اند و این مسئله باعث می‌شود که با افزایش مقدار عددی مرکز متغیر ورودی، پهنای پاسخ برآورد شده را افزایش یابد، اما اگر این روند برقرار نباشد این مسئله کارایی مدل برآزش شده را پائین می‌آورد. از این رو کائو-چیو [۷۷] طی یک روش دو مرحله‌ای به حل این مشکل پرداختند. آن‌ها با در نظر گرفتن ورودی و خروجی فازی و ضرایب دقیق، مدلی را ارائه کردند که دارای یک جمله خطای فازی بود. در مرحله اول با استفاده از روش مرکز ثقل به صورت زیر اعداد فازی را به اعداد حقیقی تبدیل نموده و ضرایب مدل را تحت روش رگرسیون کلاسیک برآورد نمودند.

$$x_{ic} = \frac{1}{3}(x_{il} + x_{im} + x_{ir}) \quad (32.1)$$

$$y_{ic} = \frac{1}{3}(y_{il} + y_{im} + y_{ir}) \quad (33.1)$$

که در آن اعداد فازی مثلی عبارتند از،

$$\tilde{x}_i = (x_{im}, x_{il}, x_{ir}) \quad \tilde{y}_i = (y_{im}, y_{il}, y_{ir}) \quad (34.1)$$

سپس براساس معیار ارائه شده توسط کیم و بیشو [۸۲] تحت می‌نیمم کردن مجموع معیار بدست آمده برای هر نمونه، برای خطای مدل پهنای چپ و راست تولید نموده و برآورد پاسخ‌ها را با فرض مثلثی بودن توابع عضویت $\tilde{x}$ و $\tilde{y}$ به صورت زیر بدست آوردند

$$\begin{aligned} Z &= \min \sum_{i=1}^n D_i \\ \text{s. t. } \quad \hat{y}_i &= b_0 + b_1 \tilde{x}_i + \tilde{e}_i \\ \tilde{e}_i &= (-\varepsilon_l, 0, \varepsilon_r) \\ (\hat{y}_{im} - \hat{y}_{il}) &\geq \varepsilon_{l(\min)}, \quad (\hat{y}_{ir} - \hat{y}_{im}) \geq \varepsilon_{r(\min)} \end{aligned} \quad (۳۵.۱)$$

که در آن داریم،

$$D_i = \frac{1}{\varphi} (y_{ir} - y_{il}) + \frac{1}{\varphi} [\varepsilon_r + \varepsilon_l + b_1 (x_{ir} - x_{il})] - 2 \times \frac{1}{\varphi} [y_{ir} - (b_0 + b_1 x_{il} - \varepsilon_l)] \lambda_i \quad (۳۶.۱)$$

و $\varepsilon_{l(\min)}$ و $\varepsilon_{r(\min)}$ به ترتیب کوچکترین پهنای چپ و راست پاسخهای مشاهده شده، و

$$\lambda_i = \frac{y_{ir} - \hat{y}_{il}}{y_{ir} - \hat{y}_{im} + \hat{y}_{im} - \hat{y}_{il}} \quad (۳۷.۱)$$

و همین‌طور b_0 و b_1 برآورد ضرایب از رگرسیون کلاسیک هستند.

در ادامه، در سال ۲۰۰۳ برای بهبود روش خود و پیوستگی بحث فازی بودن مدل و بررسی آن از طریق روش‌های فازی آن‌ها در [۷۸] با استناد به محدوده مقادیر بدست آمده تحت α - برش دلخواه، بازه‌هایی به صورت $(\varepsilon_i)_\alpha = [(\varepsilon_i)_\alpha^L, (\varepsilon_i)_\alpha^U]$ برای هر خطای فازی $\tilde{\varepsilon}_i$ ارائه نمودند که در آن $(\varepsilon_i)_\alpha^L$ و $(\varepsilon_i)_\alpha^U$ به ترتیب حد پایین و بالای این بازه هستند و داریم،

$$(\varepsilon_i)_\alpha = \begin{cases} [(Y_i)_\alpha^L - b_0 - b_1 (X_i)_\alpha^U, (Y_i)_\alpha^U - b_0 - b_1 (X_i)_\alpha^L] & b_1 \geq 0 \\ [(Y_i)_\alpha^L - b_0 - b_1 (X_i)_\alpha^L, (Y_i)_\alpha^U - b_0 - b_1 (X_i)_\alpha^U] & b_1 < 0 \end{cases} \quad (۳۸.۱)$$

و در نهایت تحت هر α - برش دلخواه تابع هدف به صورت زیر در خواهد آمد،

$$Z_\alpha = [Z_\alpha^L, Z_\alpha^U] = \left[\sum_{i=1}^n (\varepsilon_i)_\alpha^L, \sum_{i=1}^n (\varepsilon_i)_\alpha^U \right] \quad (۳۹.۱)$$

با بدست آمدن ε_i ها تحت می‌نیمم کردن Z_α ، میانگین حسابی $\tilde{\varepsilon}_i$ ها به عنوان خطای فازی مدل حاصل گردیده و برآوردهای مورد نظر تولید می‌گردند.

علاوه بر آن مسئله را برای سه حالت ورودی حقیقی و خروجی فازی، ورودی و خروجی فازی و مشاهدات فازی غیر مثلثی انجام دادند. در واقع اهمیت ویژه این مقاله آن است که برای تمام انواع مشاهدات فازی بحث شده است و همین‌طور توانایی توزیع‌پذیری بهتر ضرایب رگرسیونی را دارا است. سپس با تعریف زیر:

$$\begin{aligned} (\varepsilon_i)_\alpha^L &= \max \left[0, \min \left\{ [(\varepsilon_i)_\alpha^L]^2, (\varepsilon_i)_\alpha^L \times (\varepsilon_i)_\alpha^U, [(\varepsilon_i)_\alpha^U]^2 \right\} \right] \\ (\varepsilon_i)_\alpha^U &= \max \left\{ [(\varepsilon_i)_\alpha^L]^2, (\varepsilon_i)_\alpha^L \times (\varepsilon_i)_\alpha^U, [(\varepsilon_i)_\alpha^U]^2 \right\} \end{aligned} \quad (۴۰.۱)$$

در [۱۰۰] کائو-چیو توانستند از افزایش پهناها تحت افزایش مقدار مراکز جلوگیری نمایند. اما روش فوق مشکلات زیر را دربر داشت که راهکاری برای رفع آن‌ها ارائه گردید.

(الف) ثابت بودن پهناهای پاسخ‌های برآورد شده

از جمله کمبودهای مدل کائو-چیو آن است که برای متغیرهای ورودی غیرفازی پهنای پاسخ‌های حاصله ثابت است و نیز در حالتی که ورودی فازی باشد این نقیصه وجود دارد.

(ب) ناهمگونی پهناهای پاسخ برآورد شده و مشاهده شده

می‌دانیم که از ویژگی‌های بهینگی مدل حفظ روند پاسخ مشاهده شده براساس پاسخ‌های برآورد شده است. اما این مسئله در مدل کائو-چیو در نظر گرفته نشده است.

(ج) عدم وجود شرط مثبت بودن پهناها

در مدل کائو-چیو، تضمینی برای نامنفی بودن پهناها وجود ندارد.

مدل چن و دانگ (۲۰۰۸)

چن و دانگ [۳۱] با مطالعه مدل کائو-چیو [۷۷] به رفع نقیصه‌های الف و ب طی یک فرآیند سه مرحله‌ای پرداختند. آن‌ها ابتدا با فرض فازی بودن ضرایب رگرسیونی و براساس اصل توسیع تحت α -برش دلخواه براساس روابط موجود در رگرسیون کلاسیک، حد بالا و پایین ضرایب را ساخته تا طی آن تابع عضویت ضرایب رگرسیونی را به‌دست آورند. سپس در مرحله دوم با استفاده از روش مرکز ثقل، این ضرایب را غیرفازی نموده تا در مرحله سوم برای هر مشاهده یک جمله خطای فازی تحت یک مسئله برنامه‌ریزی، جملات خطا را برآورد نماید. در این مسیر برای برآورد خطاها، تابع هدف را اختلاف بین پاسخ‌های مشاهده شده و برآورد شده متناظر با یکدیگر مد نظر قرار دادند که در مدل کیم-بیشو [۸۲] بکار رفته بود. آن‌ها بهینگی مدل خود را در حالت‌هایی که ورودی فازی و یا غیر فازی باشد تحت مثال‌هایی نشان دادند.

مدل لو و وانگ (۲۰۰۹)

در سال ۲۰۰۹ لو و وانگ [۹۰] بر اساس دو مدل، یکی با ورودی غیرفازی و خروجی و ضرایب فازی و دیگری با فرض فازی بودن تمام عوامل مقاله خود را ارائه کردند. ایده اصلی بحث شده در [۹۰] عدم تبعیت پهناهای مشاهدات برآورد شده از روند موجود در مشاهدات مستقل است [۳۱، ۷۷، ۷۸]. آن‌ها با نظر به اینکه در مدل رگرسیونی فازی پهناها و مقادیر متغیرهای مستقل اطلاعات در دسترس ما هستند رابطه پهناها و مراکز متغیر وابسته را در دو حالت زیر فرض کردند:

الف) این مقادیر با روند موجود در پهنایها و مقادیر متغیرهای مستقل هماهنگ هستند [۷۷]، [۷۸].

ب) روند موجود در متغیر وابسته روند مجزا از روندهای موجود در متغیرهای مستقل دارند. آن‌ها برای رفع این مسئله با فرض ضرایب دقیق و محاسبه آن‌ها از روی روش رگرسیون کلاسیک، با ارائه یک جمله خطای فازی تحت مدل زیر از روی مقادیر مراکز و پهنایهای متغیرهای مستقل به برآورد پهنایها و مراکز متغیر وابسته پرداختند.

$$\begin{aligned} \min \quad & Z = \frac{1}{n} \sum_{i=1}^n D_i \\ \text{s. t.} \quad & D_i = \frac{\int \min(\mu_{\tilde{Y}_i}(x), \mu_{\hat{Y}_i}(x)) dx}{\int \max(\mu_{\tilde{Y}_i}(x), \mu_{\hat{Y}_i}(x)) dx} \\ & \hat{Y}_i = k_0 + k_1 m_{x_{i1}} + \dots + k_p m_{x_{ip}} + \tilde{E}_i \\ & \tilde{E}_i = (\circ, \alpha_{E_i}, \beta_{E_i})_{LR} \\ & \alpha_{E_i} = \sum_{j=1}^p k_{llj} \alpha_{x_{ij}} + \sum_{j=1}^p k_{lmj} m_{x_{ij}} + \sum_{j=1}^p k_{lrj} \beta_{x_{ij}} + c_l \\ & \beta_{E_i} = \sum_{j=1}^p k_{rlj} \alpha_{x_{ij}} + \sum_{j=1}^p k_{rmj} m_{x_{ij}} + \sum_{j=1}^p k_{rrj} \beta_{x_{ij}} + c_r \\ & \alpha_{E_i}, \beta_{E_i} \geq \circ, \quad i = 1, 2, \dots, n, \quad j = \circ, 1, \dots, p \end{aligned}$$

با حل مسئله فوق پارامترهای k_j ، k_{llj} ، k_{lmj} ، k_{lrj} ، k_{rlj} ، k_{rrj} ، c_l و c_r برآورد می‌شوند. این پارامترها می‌توانند خاصیت منفی شدن را داشته باشند از این رو برای داشتن روند دچار مشکل نخواهند شد.

مدل مقدس و همکاران (۲۰۱۳)

در این پایان‌نامه به عیب موجود در روش کائو-چیو [۷۸] یعنی ثابت بودن پهنایهای برآورد تحت ورودی‌های دقیق اشاره شده است. به علاوه در آن مدل شرط مثبت بودن مقادیر پهنایها لحاظ نشده است. پس آن‌ها برای رفع این مسئله و همین‌طور تعمیم مدل برای حالت‌های کلی تغییرات زیر را در مدل کائو-چیو [۷۸] اعمال کردند:

الف) بجای جمله خطای کلی، برای هر مشاهده از نمونه یک جمله خطا به صورت $E_i = (-l_i, \circ, r_i)$ ارائه گردید که اثر آن ایجاد تناسب بین پهنای پاسخ‌های مشاهده شده و برآورد شده است.

ب) قرار دادن شرط مثبت بودن پهناهای جملات خطا

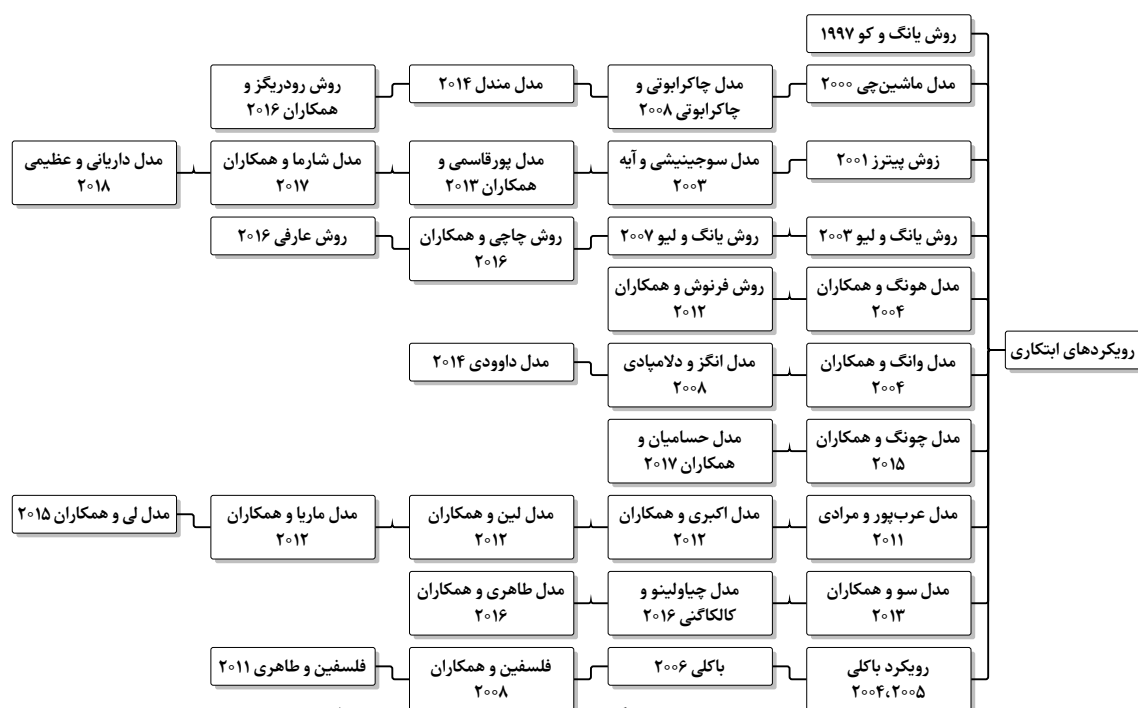
و در نهایت با ارائه مثال‌هایی برتری آن را نسبت به مدل‌های مشابه نشان دادند.

۶.۱ مدل‌های ابتکاری

در دو دهه اخیر رگرسیون فازی علاوه بر رویکردهای فوق در موارد دیگر نیز بکار گرفته شده که از آن جمله می‌توان به روش‌های مبتنی بر قواعد منطق اگر-آنگاه، خوشه بندی، روش‌های ناپارامتری و نیمه پارامتری، مدل‌های جریمه‌ای، شبکه‌های عصبی و غیره اشاره کرد (رجوع کنید به [۳]).

در این بخش برآنیم تا روند تاریخی این روش‌ها را به طور خلاصه بررسی نماییم. این بررسی را به رویکرد دسته بندی روش‌ها و تاریخ شکل‌گیری آن‌ها انجام خواهیم داد. قبل از بیان مدل‌های مختلف ابتکاری، روند تاریخی و ساختارهای مختلف این رویکرد را در قالب درخت رشد رویکرد ابتکاری به صورت زیر نمایش می‌دهیم.

درخت رشد رگرسیون فازی در رویکرد ابتکاری



شکل ۷.۱: درخت رشد رگرسیون فازی در رویکرد ابتکاری

مدل‌های خوشه‌بندی

این روش در سال ۱۹۹۷ توسط یانگ و کو [۱۴۳] پیشنهاد گردید. ایشان با توجه به این مطلب که در مواجهه با داده‌های ناهمگون و چند طیفی، سهم هر گروه در اثر گذاری یکسان نیست، این ایده را داشت که بخش‌های مختلف می‌بایست با سهم اثرهای متفاوت حضور داشته باشند. از سوی دیگر، استفاده از روش خوشه بندی ما را به الگوهای بهتری راهنمایی خواهد کرد بنابراین از درجه عضویت به عنوان وزن برای هر مشاهده درخوشه‌های مختلف استفاده کرده و مدل خود را ارائه کردند. در ادامه و در سال ۲۰۰۱، پیترز [۱۰۳] از روش خوشه‌بندی- رگرسیونی به عنوان یک روش کاربردی در تحلیل داده‌های اقتصادی بهره برد. همچنین یانگ و لین [۱۴۴] نیز برای مدلی با ورودی، ضرایب و خروجی فازی، روش کمترین توان‌های دوم با خوشه بندی با فرض اینکه ضرب دو عدد فازی تقریباً یک عدد فازی است مطرح کرد. در ادامه و در سال ۲۰۰۳، یانگ و لیو [۱۴۵] این روش را برای کم اثر کردن نقاط دورافتاده در مدل رگرسیونی فازی مورد استفاده قرار دادند. علاوه بر آن، استفاده از این روش در یک مسئله کاربردی و عملی، توسط چاچی و همکاران [۲۷] در سال ۲۰۱۶ مورد استفاده قرار گرفت. در این مقاله، این روش برای برآورد بار معلق در رودخانه (در بحث هیدرولوژی) استفاده گردید. در همان سال این رویکرد توسط عارفی [۱۴] برای حالتی که مشاهدات متغیر پاسخ در آن‌ها به صورت اعداد بازه‌ای مقدار باشند توسعه داده شد.

مدل‌های پارامتری و ناپارامتری

در مباحث آماری زمانی که توزیع مشاهدات مشخص باشد می‌توان از روش‌های پارامتری بهره کافی برد زیرا در آنجا از شاخص‌هایی که به توزیع وابسته هستند و علاوه بر آن تغییرات توزیع را بهتر از شاخص‌های دیگر در خود نشان می‌دهند مانند میانگین و واریانس استفاده می‌گردد. اما زمانی که توزیع مشاهدات نامشخص باشد و همچنین شرایط برای استفاده از قضایای مهم مانند قضیه حد مرکزی فراهم نباشد، مجبور به استفاده از روش‌های ناپارامتری و یا نیمه پارامتری خواهیم بود.

در این مسیر بکارگیری روش‌های جریمه‌ای به عنوان مدل‌های پارامتری و ناپارامتری که مدل‌های ریج و لاسو از جمله آن‌ها به شمار می‌رود، در آمار کلاسیک هم به لحاظ نظری و هم به لحاظ کاربردی در طیف گسترده‌ای مورد استقبال و استفاده قرار گرفته است. این مسئله در مورد مباحث فازی نیز برقرار است و در سال‌های اخیر محققان در این زمینه تحقیقات متعددی را به انجام رسانیده‌اند.

در سال ۲۰۰۴، هونگ و همکاران [۷۰] مدلی را پیشنهاد دادند که در آن رگرسیون ریج مورد بحث قرار گرفت. در این مدل مشاهدات ورودی غیر فازی بودند و خروجی و ضرایب این مدل نیز در دو حالت مثلثی متقارن و نرمال فازی بررسی گردید و تحت مثال‌هایی برتری

مدل را بیان کردند. در سال ۲۰۰۷ وانگ و همکاران [۱۳۸] نیز با استفاده از روش کمترین توان‌های دوم، یک مدل رگرسیونی ناپارامتری فازی را ارائه کردند. انگز و دلامپای [۱۲] نیز با استفاده از مجموعه‌های فازی در سال ۲۰۰۸ یک مدل رگرسیونی ناپارامتری بیزی را پیشنهاد دادند. همچنین در سال ۲۰۱۲ با استفاده از فاصله دیاموند، فرنوش و همکاران [۵۲] به برآورد پارامترهای رگرسیون ریج در محیط فازی پرداختند. داوودی [۱] نیز رگرسیون ناپارامتری فازی را با بکارگیری موجک‌ها مورد مطالعه قرار داد.

در سال ۲۰۱۷ حسامیان و همکاران [۶۳] با استفاده از متر قدرمطلق ارائه شده در [۸۱] به برآورد پارامترهای یک مدل رگرسیونی فازی نیمه پارامتری پرداختند. آن‌ها با در نظر گرفتن ورودی و خروجی فازی به صورت یک عدد فازی LR مدل خود را ارائه و تحت مثال‌هایی واقعی و شبیه‌سازی شده برتری روش خود را نشان دادند. همچنین در [۶۶] رگرسیون لاسو به عنوان یک مدل جریمه‌ای مورد استفاده قرار گرفت. در این مدل آن‌ها با جریمه کردن مراکز و پهنای ضرایب به طور جداگانه، که با ورودی‌های غیرفازی بوده است، تحت داده‌های شبیه‌سازی شده و مثال‌های واقعی توانایی مدل لاسو را در مقابل مدل‌های دیگر به نمایش گذاشتند.

در سال ۲۰۱۹ نیز اکبری و حسامیان [۹] مدل رگرسیون فازی جریمه‌ای الاستیک نت را (که ترکیبی از مدل‌های ریج و لاسو است) با ورودی و خروجی فازی ارائه دادند. آن‌ها در این مقاله با استفاده از یک مدل نیمه پارامتری و مدل‌های جریمه‌ای به دنبال دستیابی به متغیرهای توضیحی مهم و انتخاب متغیر از این طریق پرداختند. آن‌ها عملکرد خوب مدل خود را با مثال‌هایی عددی و شبیه‌سازی ارائه داده و برتری آن را در مقایسه با چند مدل رگرسیونی ریج فازی متداول نشان دادند.

در سال ۲۰۱۹ نیز ربیعی و همکاران [۱۱۰]، تحت یک مدل تمام فازی، با استفاده از مدل ارائه شده توسط حسن پور [۶۱] مدل رگرسیونی ریج را ارائه و با مثال‌هایی برتری آن را نسبت به مدل اولیه نشان دادند.

روش بوت استرپ

یکی از روش‌های مورد استفاده در مواقعی که حجم نمونه پایین است، استفاده از روش‌های بازنمونه‌گیری مانند بوت استرپ است. این روش توسط افرون در سال ۱۹۷۹ [۴۷] ابداع گردید و سپس در سال ۱۹۹۴ [۴۸] به همراه تیبشیرانی کاربردهای مختلف این روش را خصوصاً در رگرسیون به تفسیر توضیح و تشریح نمود. در مباحث فازی نیز این روش مورد استفاده قرار گرفت. در سال ۲۰۱۱ عرب پور و مرادی [۴] با استفاده از روش باز نمونه‌گیری بوت استرپ در رگرسیون کمترین توان‌های دوم در نمونه‌های با حجم کم برتری این روش را در مقاله‌هایی عددی نسبت به دیگر روش‌ها نشان دادند. در سال ۲۰۱۲ اکبری و همکاران [۸] نیز این رویکرد را مورد استفاده قرار دادند. در همان سال لین و همکاران [۸۷] این رویکرد را در یک مثال

عملی برای تحلیل داده‌های اقتصادی بکار بردند. در این مسیر فاصله اطمینان بوت استریپی برای پارامترهای مدل‌های رگرسیونی فازی نیز مورد استفاده محققین قرار گرفت که از جمله آن‌ها می‌توان به [۵۳] اشاره کرد.

بهره‌گیری از رویکرد بوت استرپ در استنباط‌های آمار فازی نیز مورد استقبال قرار گرفت و در سال ۲۰۱۵، لی و همکاران [۸۵] با استفاده از فاصله اطمینان‌های بوت استریپی تحت آلفا-برش‌های مختلف، استنباط‌هایی در مورد پارامترهای فازی برآورد شده در مدل رگرسیونی فازی انجام دادند.

مدل‌های استوار

این مقوله نیز بسیار مورد توجه محققین قرار دارد چرا که در عمل نقاط دور افتاده مدل را تحت تاثیر خود قرار داده و پیش‌بینی مدل را به چالش می‌کشند. از این رو قدرتمند بودن مدل رگرسیونی درمقابل اثرات نقاط پرت از مزیت‌های آن به شمار می‌آید. در رگرسیون فازی نیز محققین به این مسئله مهم پرداخته و گه‌گاه به عنوان یک مزیت در مدل‌های خود این خاصیت را نشان داده‌اند. برای مدل‌هایی از این دست می‌توان به مقالات کلکین نما و طاهری [۸۱] و دورسو و ماساری [۴۶] اشاره کرد.

این دیدگاه در سال ۲۰۱۶ توسط چاچی و روزبه [۲۷] به صورت کاربردی برای یک مثال کاربردی در داده‌های حمل بار مورد استفاده قرار گرفت.

رویکرد باکلی

در سال ۲۰۰۴ در [۱۹] و سال ۲۰۰۵ در [۲۰، ۲۱]، باکلی دیدگاه جدیدی را ارائه کرد. او با توجه به ساختار بازه اطمینان کلاسیک که بر اساس درصد اطمینان‌های مختلف بازه‌های تو در تویی را در محدوده صفر تا ۱۰۰ تولید می‌کنند، آن را متناظر با ساختار آلفا-برش در اعداد فازی در نظر گرفت. او این دیدگاه را در بخش‌های مختلف احتمال و استنباط کلاسیک توسعه داد و این مطالب را در کتابی به طور مفصل با ارائه مثال‌هایی تشریح کرد (رجوع کنید به [۲۲]). در ادامه فلسفی و همکاران [۴۹]، در برآورد پارامترها از این شیوه استفاده کردند. در سال ۲۰۱۱، فلسفی و طاهری [۵۰] با بررسی این دیدگاه در مدل بندی رگرسیونی، استفاده از آن را برای برآورد میانگین تایید کردند. اما برای بکارگیری آن در مورد برآورد واریانس، با ایجاد اصلاحاتی آن را بهبود دادند. این دیدگاه در مباحث کنترل کیفیت بسیار مورد استقبال واقع شده و بکار رفته است که از جمله می‌توان از [۱۰۱، ۱۰۲] نام برد.

رویکردهای دیگر

به غیر از مدل‌های اشاره شده رویکردهای دیگری نیز در رگرسیون، خصوصاً رگرسیون فازی مورد استفاده و بهره‌برداری قرار گرفته‌اند که در زیر به اختصار به آن‌ها اشاره می‌گردد. یکی از این ساختارها استفاده از سیستم‌های مبتنی بر قواعد آنگاه-اگر است. این مدل‌ها منجر به بروز یک مدل خاص و مشهود نمی‌گردد اما توانایی پیش‌بینی متغیر پاسخ را داراست. مقالات [۲۹، ۵] از این دست می‌باشند. در سال‌های اخیر نیز این روش توسط مندل [۹۱] و رودریگز و همکاران [۱۱۴] مورد استفاده قرار گرفته است.

گروه دیگر از روش‌های ابتکاری را می‌توان مدل‌های شبکه‌های عصبی [۱۲۱، ۳۹] و مدل‌های SVM [۱۳۶]، یادگیری ماشین [۸۸]، الگوریتم‌های EM [۱۲۵]، AHP [۱۰۶]، و ماکسیمم آنترופی [۳۳] در نظر گرفت که داده‌های فازی را برای مدل‌های خود مورد استفاده قرار داده‌اند.

فصل ۲

برآورد ضرایب مدل رگرسیون فازی به روش جک نایف بعد از بوت استرپ

۱.۲ مقدمه

تقریباً هدف اصلی در تمام زمینه‌های علمی مانند زیست‌شناسی، اقتصاد، مهندسی و پزشکی تبیین و پیش‌بینی خروجی سیستم با برخی از متغیرهای ورودی اکتشافی است. یکی از راه‌های دستیابی به این هدف استفاده از مدل رگرسیون به عنوان یک رابطه ریاضی بین متغیر خروجی یعنی پاسخ و متغیرهای ورودی یعنی پیش‌بینی‌کننده‌هاست. روش‌های معمول شامل کمترین توانهای دوم (LS) و یا درست‌نمایی ماکزیمم (ML)، هنگامی که نمونه‌های کمی وجود دارد، نتایج قابل اعتماد و رضایت‌بخشی ارائه نمی‌دهند. در واقع بخاطر عدم حصول به توزیع داده‌ها، استنباط‌های مورد استفاده در این مسیر دچار نقص و اشکال خواهند شد. از این‌رو برای بررسی، تحلیل و مدل‌بندی اینگونه متغیرها از روش‌های ناپارامتری بهره خواهیم برد. در این بین روش‌هایی با نام «روش‌های باز نمونه‌گیری» ما را در افزایش نمونه مشاهده شده، تخمین اریبی و خطای معیار و دیگر مسائل استنباطی کمک خواهند کرد.

۲.۲ مثال‌های انگیزشی

در این بخش مجموعه داده‌های واقعی را که در فصل‌های بعد نیز برای بررسی و نمایش عملکرد روش‌ها و تکنیک‌های پیشنهادی مورد استفاده قرار گرفته اند، در قالب مثال‌های انگیزشی آورده‌ایم.

مثال ۱.۲.۲. مجموعه داده اول، داده‌های واقعی جدول ۱.۲ برای اولین بار توسط کیم و بیشو [۸۲] استفاده شده است. این مجموعه داده‌ها مربوط به مطالعه‌ای درباره زمان پاسخ تشخیص خدمه اتاق کنترل نیروگاه هسته‌ای به یک واقعه غیر عادی است. زمان پاسخ تشخیص خدمه در شرایط غیر عادی معمولاً تحت تأثیر عوامل شکل دهنده عملکرد مانند سطح استرس، پتانسیل تشخیص اشتباه و غیره قرار می‌گیرد. فرض بر این است که لگاریتم زمان پاسخ به طور نرمال توزیع می‌شود. با این فرض، پهنای چپ و راست اعداد فازی با ۹۵ درصد حد پایین و بالا از توزیع‌های غیر نرمال داده می‌شوند. در اینجا، یک مدل خطی فازی بین زمان پاسخ شناختی و تجربه خدمه در داخل اتاق کنترل (بر حسب سال) و تجربه خارج از اتاق کنترل (بر حسب سال) و تحصیلات (بر حسب سال) فرض می‌شود، که در آن، x_1 ، x_2 و x_3 به ترتیب تجربه خدمه در داخل اتاق کنترل، تجربه خارج از اتاق کنترل و آموزش را نشان می‌دهد.

جدول ۱.۲: مجموعه داده‌های زمان تشخیص خدمه اتاق کنترل نیروگاه هسته‌ای بر یک واقعه غیرعادی [۸۲].

ردیف	x_1	x_2	x_3	$\tilde{Y}_i$
۱	۲/۰۰	۰/۰۰	۱۵/۲۵	$(۵/۸۳, ۳/۵۶)_T$
۲	۰/۰۰	۵/۰۰	۱۴/۱۳	$(۰/۸۵, ۰/۵۲)_T$
۳	۱/۱۳	۱/۵۰	۱۱/۶۳	$(۱۳/۹۳, ۸/۵۰)_T$
۴	۲/۰۰	۱/۲۵	۱۳/۶۳	$(۴/۰۰, ۲/۴۴)_T$
۵	۲/۱۹	۳/۷۵	۱۴/۷۵	$(۱/۵۸, ۰/۹۶)_T$
۶	۰/۲۵	۳/۵۰	۱۳/۷۵	$(۱/۵۸, ۰/۹۶)_T$
۷	۰/۷۵	۵/۲۵	۱۵/۲۵	$(۸/۱۸, ۴/۹۹)_T$
۸	۴/۲۵	۲/۰۰	۱۳/۵۰	$(۱/۸۵, ۱/۱۳)_T$

مثال ۲.۲.۲. مجموعه داده دوم، داده‌های واقعی جدول ۲.۲ در یک آزمایش واقعی از ارزیابی تأثیر سه متغیر بر مزه پنیر تولید شده است. در داده‌ها، اسید استیک، سولفید هیدروژن و اسید لاکتیک به ترتیب سه متغیر ورودی غیرفازی هستند و متغیر خروجی نیز به عنوان طعم و مزه پنیر که توسط یک متخصص ارزیابی می‌شود. برای تولید پاسخ مثلی فازی، ما ۱۵ درصد نقاط مرکزی را به عنوان پهنای در نظر می‌گیریم.

جدول ۲.۲: مجموعه داده دوم.

ردیف	اسید استیک	سولفید هیدروژن	اسید لاکتیک	طعم و مزه پنییر
۱	۴/۵۴۳	۳/۱۳۵	۰/۸۶	$(۱۲/۳, ۱/۸۴۵)_T$
۲	۵/۱۵۹	۵/۰۴۳	۱/۵۳	$(۲۰/۹, ۳/۱۳۵)_T$
۳	۵/۳۶۶	۵/۴۳۸	۱/۵۷	$(۳۹/۰, ۵/۸۵۰)_T$
۴	۵/۷۵۹	۷/۴۹۶	۱/۸۱	$(۴۷/۹, ۰/۱۸۵)_T$
۵	۴/۶۶۳	۳/۸۰۷	۰/۹۹	$(۵/۶, ۰/۸۴۰)_T$
۶	۵/۶۹۷	۷/۶۰۱	۱/۰۹	$(۲۵/۹, ۳/۸۸۵)_T$
۷	۵/۸۹۲	۸/۷۲۶	۱/۲۹	$(۳۷/۳, ۵/۵۹۵)_T$
۸	۶/۰۷۸	۷/۹۶۶	۱/۷۸	$(۲۱/۹, ۳/۲۸۵)_T$
۹	۴/۸۹۸	۳/۸۵	۱/۲۹	$(۱۸/۱, ۲/۷۱۵)_T$
۱۰	۵/۲۴۲	۴/۱۷۶	۱/۵۸	$(۲۱/۰, ۳/۱۵۰)_T$
۱۱	۵/۷۴	۶/۱۴۲	۱/۶۸	$(۳۴/۹, ۵/۲۳۵)_T$
۱۲	۶/۴۴۶	۷/۹۰۸	۱/۹	$(۵۷/۲, ۸/۵۸۰)_T$
۱۳	۴/۴۷۷	۲/۹۹۶	۱/۰۶	$(۰/۷۰, ۰/۱۰۵)_T$
۱۴	۵/۲۳۶	۴/۹۴۲	۱/۳۰	$(۲۵/۹, ۳/۸۸۵)_T$
۱۵	۶/۱۵۱	۶/۷۵۲	۱/۵۲	$(۵۴/۹, ۸/۲۳۵)_T$

مثال ۳.۲.۲. داده‌های ساختگی ساکاوا و یانو [۱۱۶]: این داده‌ها که به عنوان مجموعه داده‌های سوم در نظر گرفته شده است، در [۱۱۶] برای اولین بار توسط ساکاوا و یانو بکارگرفته شد.

جدول ۳.۲: مجموعه داده‌های ساختگی ساکاوا و یانو [۱۱۶].

ردیف	$\tilde{X}_i$	$\tilde{Y}_i$
۱	$(۲/۰۰, ۰/۵۰)_T$	$(۴/۰۰, ۰/۵۰)_T$
۲	$(۳/۵۰, ۰/۵۰)_T$	$(۵/۵۰, ۰/۵۰)_T$
۳	$(۵/۵۰, ۱/۰۰)_T$	$(۷/۵۰, ۱/۰۰)_T$
۴	$(۷/۰۰, ۰/۵۰)_T$	$(۶/۵۰, ۰/۵۰)_T$
۵	$(۸/۵۰, ۰/۵۰)_T$	$(۸/۵۰, ۰/۵۰)_T$
۶	$(۱۰/۵۰, ۱/۰۰)_T$	$(۸/۰۰, ۱/۰۰)_T$
۷	$(۱۱/۰۰, ۰/۵۰)_T$	$(۱۰/۵۰, ۰/۵۰)_T$
۸	$(۱۲/۵۰, ۰/۵۰)_T$	$(۹/۵۰, ۰/۵۰)_T$

مثال ۴.۲.۲. داده‌های قیمت خانه‌های شانگهای (مجموعه داده چهارم)، این داده‌ها اولین بار، ابتدا توسط [۱۵۱] معرفی شدند. آن‌ها قیمت خرید قابل قبول خانه را به عنوان یک عدد فازی مثلثی نامتقارن با شش متغیر توضیحی غیرفازی در نظر گرفتند که در آن، x_1 : اندازه مسکن. x_2 : نرخ بهره وام مسکن ؛ x_3 : مالیات بر املاک و مستقالات x_4 : نسبت پیش پرداخت x_5 : درآمد سالانه خانوار و در نهایت x_6 : جمعیت خانواده را به عنوان متغیرهای ورودی در نظر گرفتند.

جدول ۴.۲: مجموعه داده‌های قیمت خانه‌های شانگهای [۱۵۱].

$\tilde{y}_i$	x_6	x_5	x_4	x_3	x_2	x_1	i	$\tilde{y}_i$	x_6	x_5	x_4	x_3	x_2	x_1	i
$(650, 280, 300)_T$	3	36	35	0/6	4/9	100	75	$(300, 120, 80)_T$	4	27	20	0/6	5/9	80	1
$(450, 100, 50)_T$	3	20	35	0/6	4/9	80	76	$(470, 140, 140)_T$	4	30	35	0	4/9	90	2
$(250, 80, 80)_T$	4	20	20	0	5/9	80	77	$(230, 100, 100)_T$	3	35	30	0/4	6/55	70	3
$(580, 130, 170)_T$	5	35	30	0	6/15	100	78	$(750, 220, 180)_T$	3	40	30	0/6	6/8	120	4
$(470, 170, 130)_T$	3	37	30	0/6	6/8	80	79	$(600, 200, 200)_T$	4	33	35	0	4/9	100	5
$(500, 180, 120)_T$	2	25	20	0/4	5/9	100	80	$(670, 240, 150)_T$	5	40	30	0	6/55	100	6
$(350, 190, 30)_T$	3	20	30	0	6/55	90	81	$(220, 100, 50)_T$	5	18	20	0/4	5/9	80	7
$(670, 250, 60)_T$	2	37	30	0	6/8	100	82	$(630, 180, 170)_T$	5	45	30	0/6	6/8	90	8
$(470, 170, 110)_T$	3	35	20	0/6	6/8	80	83	$(200, 100, 50)_T$	4	12	30	0/6	6/15	90	9
$(530, 150, 120)_T$	2	40	20	0/6	6/15	90	84	$(550, 170, 80)_T$	3	30	30	0	6/15	100	10
$(430, 150, 70)_T$	3	30	30	0	6/55	80	85	$(350, 130, 50)_T$	3	25	35	0	4/9	80	11
$(170, 70, 60)_T$	3	28	30	0/4	6/15	60	86	$(330, 130, 90)_T$	3	22	20	0/6	5/9	90	12
$(180, 60, 70)_T$	3	22	35	0/6	4/9	70	87	$(1000, 300, 270)_T$	3	65	30	0/6	6/8	100	13
$(650, 170, 270)_T$	3	43	30	0/6	4/7	100	88	$(180, 80, 70)_T$	4	23	30	0/6	6/15	70	14
$(470, 170, 110)_T$	2	30	35	0/6	4/9	90	89	$(370, 170, 100)_T$	4	30	30	0/6	6/55	80	15
$(270, 120, 50)_T$	4	27	35	0/6	4/9	70	90	$(200, 80, 70)_T$	5	23	30	0/4	4/9	70	16
$(830, 250, 220)_T$	3	47	20	0/6	5/9	110	91	$(620, 170, 110)_T$	5	34	30	0	6/15	100	17
$(250, 80, 70)_T$	2	20	30	0	4/7	80	92	$(220, 100, 30)_T$	4	20	20	0/6	5/9	80	18
$(620, 220, 80)_T$	4	35	30	0/4	6/8	100	93	$(200, 100, 70)_T$	2	18	20	0/6	5/9	80	19
$(430, 200, 120)_T$	4	32	40	0/4	4/9	80	94	$(680, 180, 190)_T$	4	42	30	0	6/55	100	20
$(780, 300, 20)_T$	3	40	30	0	6/15	100	95	$(420, 190, 60)_T$	3	30	35	0/6	4/9	80	21
$(480, 150, 40)_T$	3	35	30	0	6/15	80	96	$(470, 220, 30)_T$	3	20	30	0	6/15	100	22
$(420, 190, 10)_T$	4	27	30	0/6	5/9	90	97	$(280, 130, 50)_T$	4	23	20	0/6	6/15	80	23
$(170, 70, 80)_T$	4	23	30	0/6	5/9	70	98	$(680, 180, 140)_T$	5	35	30	0	6/8	110	24
$(320, 100, 60)_T$	2	25	30	0	5/9	80	99	$(430, 150, 40)_T$	3	33	30	0/6	3/25	80	25
$(220, 90, 50)_T$	3	18	30	0/4	3/25	80	100	$(720, 190, 180)_T$	4	47	30	0	6/15	100	26
$(350, 150, 20)_T$	4	30	30	0	6/15	70	101	$(1130, 400, 90)_T$	5	57	30	0/4	4/7	120	27
$(1000, 320, 150)_T$	3	50	30	0	6/8	120	102	$(380, 110, 100)_T$	3	23	30	0	5/9	90	28
$(530, 180, 50)_T$	4	33	30	0	6/15	90	103	$(630, 200, 70)_T$	4	35	30	0	5/9	100	29
$(220, 90, 50)_T$	4	23	30	0	6/8	70	104	$(350, 80, 150)_T$	3	25	30	0	4/7	90	30
$(580, 200, 150)_T$	3	35	30	0/6	6/8	100	105	$(600, 200, 180)_T$	4	37	30	0	6/8	100	31
$(570, 170, 80)_T$	3	35	35	0	4/9	90	106	$(370, 150, 50)_T$	3	28	20	0/4	6/15	80	32
$(200, 80, 80)_T$	2	25	30	0/6	6/15	70	107	$(320, 120, 60)_T$	3	20	30	0	6/55	90	33
$(300, 100, 30)_T$	3	25	30	0/4	6/15	80	108	$(500, 150, 150)_T$	4	32	30	0/4	6/55	100	34
$(570, 170, 80)_T$	3	32	20	0	6/55	100	109	$(550, 200, 150)_T$	4	35	30	0	6/55	90	35

$\tilde{y}_i$	x_6	x_5	x_4	x_3	x_2	x_1	i	$\tilde{y}_i$	x_6	x_5	x_4	x_3	x_2	x_1	i
$(520, 190, 50)_T$	۳	۳۰	۲۰	۰	۵/۹	۹۰	۱۱۰	$(270, 140, 80)_T$	۴	۱۷	۲۰	۰/۶	۵/۹	۹۰	۳۶
$(1000, 270, 50)_T$	۲	۴۷	۲۰	۰	۶/۱۵	۱۲۰	۱۱۱	$(630, 230, 120)_T$	۴	۳۵	۳۰	۰	۶/۱۵	۱۰۰	۳۷
$(800, 230, 120)_T$	۴	۵۰	۳۰	۰/۴	۶/۵۵	۱۰۰	۱۱۲	$(320, 120, 60)_T$	۲	۲۳	۳۰	۰	۵/۹	۸۰	۳۸
$(300, 120, 80)_T$	۳	۲۰	۳۰	۰/۴	۶/۵۵	۹۰	۱۱۳	$(150, 50, 30)_T$	۴	۱۵	۲۰	۰/۶	۶/۸	۸۰	۳۹
$(250, 100, 70)_T$	۲	۳۰	۳۰	۰	۶/۱۵	۶۰	۱۱۴	$(350, 120, 30)_T$	۳	۲۰	۳۰	۰	۶/۱۵	۹۰	۴۰
$(250, 80, 80)_T$	۳	۲۷	۶۰	۰	۵/۹	۷۰	۱۱۵	$(320, 120, 100)_T$	۳	۲۷	۳۵	۰/۶	۴/۹	۸۰	۴۱
$(320, 70, 50)_T$	۳	۳۲	۳۰	۰	۶/۵۵	۷۰	۱۱۶	$(300, 100, 30)_T$	۳	۲۲	۳۰	۰	۶/۱۵	۸۰	۴۲
$(450, 130, 80)_T$	۳	۳۵	۴۰	۰	۴/۹	۸۰	۱۱۷	$(670, 190, 150)_T$	۴	۴۳	۳۰	۰	۶/۸	۱۰۰	۴۳
$(500, 150, 120)_T$	۵	۳۵	۲۰	۰	۵/۹	۸۰	۱۱۸	$(630, 200, 120)_T$	۵	۳۵	۳۵	۰	۴/۹	۱۰۰	۴۴
$(380, 130, 90)_T$	۲	۲۳	۳۰	۰	۶/۱۵	۹۰	۱۱۹	$(170, 50, 60)_T$	۴	۲۷	۲۰	۰/۶	۶/۵۵	۶۰	۴۵
$(620, 250, 200)_T$	۴	۳۵	۲۰	۰/۶	۶/۵۵	۱۰۰	۱۲۰	$(970, 240, 200)_T$	۳	۵۰	۳۵	۰	۴/۹	۱۲۰	۴۶
$(370, 120, 60)_T$	۳	۲۸	۳۰	۰	۵/۹	۸۰	۱۲۱	$(180, 60, 70)_T$	۳	۲۵	۳۰	۰/۶	۶/۱۵	۷۰	۴۷
$(780, 260, 170)_T$	۳	۴۸	۳۰	۰/۶	۴/۹	۱۰۰	۱۲۲	$(530, 230, 70)_T$	۳	۲۸	۳۰	۰/۴	۵/۹	۱۰۰	۴۸
$(530, 250, 50)_T$	۴	۳۷	۳۰	۰/۶	۶/۵۵	۸۰	۱۲۳	$(220, 100, 80)_T$	۳	۲۰	۲۰	۰/۶	۶/۸	۸۰	۴۹
$(180, 60, 40)_T$	۴	۲۷	۳۰	۰/۶	۵/۹	۶۰	۱۲۴	$(620, 250, 100)_T$	۴	۳۷	۳۰	۰/۴	۶/۱۵	۱۰۰	۵۰
$(600, 170, 80)_T$	۳	۴۷	۲۰	۰/۶	۶/۱۵	۸۰	۱۲۵	$(430, 150, 90)_T$	۴	۳۰	۲۰	۰	۶/۱۵	۸۰	۵۱
$(500, 100, 100)_T$	۳	۳۰	۳۰	۰	۶/۱۵	۸۰	۱۲۶	$(420, 120, 100)_T$	۳	۳۰	۳۰	۰/۴	۶/۵۵	۹۰	۵۲
$(400, 150, 70)_T$	۴	۳۰	۳۰	۰	۶/۱۵	۸۰	۱۲۷	$(250, 120, 50)_T$	۳	۲۰	۲۰	۰/۶	۵/۹	۸۰	۵۳
$(230, 80, 40)_T$	۳	۲۳	۳۰	۰	۵/۹	۷۰	۱۲۸	$(180, 60, 120)_T$	۲	۲۵	۳۵	۰/۶	۴/۹	۷۰	۵۴
$(170, 70, 60)_T$	۳	۲۳	۳۰	۰/۶	۶/۱۵	۷۰	۱۲۹	$(400, 120, 80)_T$	۴	۳۰	۲۰	۰	۵/۹	۸۰	۵۵
$(720, 240, 100)_T$	۴	۴۰	۶۰	۰	۵/۹	۱۰۰	۱۳۰	$(500, 180, 180)_T$	۳	۳۵	۳۰	۰/۶	۶/۵۵	۹۰	۵۶
$(380, 110, 100)_T$	۲	۳۳	۲۰	۰	۶/۱۵	۷۰	۱۳۱	$(400, 130, 100)_T$	۴	۳۰	۲۰	۰	۶/۸	۸۰	۵۷
$(480, 200, 90)_T$	۳	۳۴	۳۰	۰/۶	۶/۸	۹۰	۱۳۲	$(180, 80, 50)_T$	۳	۲۳	۲۰	۰/۶	۵/۹	۷۰	۵۸
$(380, 110, 90)_T$	۴	۳۰	۲۰	۰	۶/۵۵	۸۰	۱۳۳	$(550, 180, 150)_T$	۴	۳۰	۳۰	۰	۶/۱۵	۱۰۰	۵۹
$(400, 120, 100)_T$	۴	۲۷	۳۰	۰	۵/۹	۹۰	۱۳۴	$(450, 230, 50)_T$	۴	۳۳	۲۰	۰/۶	۶/۵۵	۸۰	۶۰
$(280, 100, 70)_T$	۴	۳۲	۳۰	۰	۶/۱۵	۸۰	۱۳۵	$(320, 90, 80)_T$	۴	۲۷	۳۵	۰	۴/۹	۸۰	۶۱
$(630, 180, 100)_T$	۳	۳۷	۳۰	۰/۴	۳/۲۵	۱۰۰	۱۳۶	$(380, 100, 90)_T$	۳	۲۵	۳۰	۰	۳/۲۵	۹۰	۶۲
$(420, 140, 100)_T$	۳	۳۵	۳۰	۰/۴	۶/۵۵	۸۰	۱۳۷	$(250, 100, 70)_T$	۲	۳۰	۳۰	۰/۴	۶/۵۵	۶۰	۶۳
$(520, 170, 110)_T$	۳	۳۰	۳۰	۰/۴	۶/۵۵	۱۰۰	۱۳۸	$(350, 120, 120)_T$	۴	۲۷	۳۵	۰	۴/۹	۸۰	۶۴
$(1020, 370, 110)_T$	۳	۴۷	۳۰	۰	۶/۱۵	۱۲۰	۱۳۹	$(1520, 470, 280)_T$	۵	۷۵	۳۰	۰/۶	۶/۸	۱۵۰	۶۵
$(200, 80, 80)_T$	۲	۲۵	۳۰	۰/۶	۵/۹	۷۰	۱۴۰	$(270, 100, 30)_T$	۴	۲۰	۲۰	۰	۵/۹	۸۰	۶۶
$(400, 100, 20)_T$	۳	۳۵	۳۰	۰	۵/۹	۷۰	۱۴۱	$(170, 70, 60)_T$	۴	۲۳	۳۰	۰/۶	۶/۱۵	۷۰	۶۷
$(320, 120, 110)_T$	۳	۲۰	۳۰	۰	۶/۱۵	۹۰	۱۴۲	$(300, 100, 70)_T$	۳	۲۵	۳۰	۰/۴	۵/۹	۸۰	۶۸
$(520, 170, 160)_T$	۲	۲۸	۳۰	۰	۶/۸	۱۰۰	۱۴۳	$(450, 200, 120)_T$	۲	۳۳	۲۰	۰/۶	۵/۹	۸۰	۶۹
$(180, 80, 50)_T$	۳	۱۷	۳۰	۰/۶	۶/۸	۸۰	۱۴۴	$(370, 140, 100)_T$	۳	۲۸	۳۵	۰	۴/۹	۸۰	۷۰
$(250, 70, 70)_T$	۳	۲۳	۳۰	۰/۴	۶/۱۵	۸۰	۱۴۵	$(200, 70, 50)_T$	۲	۲۳	۳۰	۰/۴	۶/۵۵	۷۰	۷۱
$(370, 100, 100)_T$	۳	۲۳	۳۵	۰	۴/۹	۹۰	۱۴۶	$(450, 120, 220)_T$	۴	۳۷	۲۰	۰	۶/۸	۸۰	۷۲
$(480, 180, 100)_T$	۴	۳۰	۳۰	۰	۶/۱۵	۹۰	۱۴۷	$(270, 120, 50)_T$	۴	۳۰	۳۰	۰/۶	۶/۱۵	۷۰	۷۳
								$(270, 120, 50)_T$	۳	۲۲	۳۰	۰/۴	۶/۵۵	۸۰	۷۴

۳.۲ بوت استرپ در رگرسیون فازی

روش بوت استرپ در سال ۱۹۷۹ توسط افرون^۱ [۴۷] ابداع گردید و در سال‌های بعد آن را توسعه داد. روش‌های بوت استرپ^۲ کلاسی از روش‌های ناپارامتری هستند که به برآورد توزیع یک جامعه با استفاده از باز نمونه‌گیری می‌پردازد. این روش می‌تواند در دو حالت با ساختارهای باجایگذاری و بدون جایگذاری نمونه‌ها انجام بگیرد. کاربرد این روش در رگرسیون به دو شکل انجام می‌پذیرد. در حالت اول روی ماتریس داده‌های ورودی و جایگزینی هر سطر از آن کار می‌شود و در حالت دوم جابجایی خطاهای بوت استرپ شده حاصل از برآورد اولیه تحت نمونه‌های اولیه را خواهیم داشت.

از سوی دیگر رگرسیون فازی به عنوان یک روش در مدل‌بندی متغیرهای دارای ابهام و یا زمانی که شرایط اولیه توزیعی در مورد خطاهای مدل برقرار نیست کارایی بسیاری دارد. علاوه بر این، کم بودن تعداد نمونه هم در کاهش اعتماد به تخمین پارامترها و هم استنباط‌های مربوط به مدل موثر است. از این رو محققین در ادامه توسعه و بهبود روش‌ها و مدل‌های فازی، استفاده از روش‌های باز نمونه‌گیری را مفید دانستند. بکارگیری این روش‌ها شامل برآورد پارامترها [۴] و تهیه فواصل اطمینان [۵۳، ۸۵، ۸۷] به روش بوت استرپ بوده است. در سال ۱۳۹۳ عرب‌پور و مرادی [۴] با فرض اینکه مدل رگرسیونی فازی به صورت

$$\tilde{y}_i = \sum_{j=0}^n \tilde{\beta}_j \tilde{x}_{ij} + \tilde{\epsilon}_i, \quad i = 1, \dots, n \quad (1.2)$$

باشد، برای برآورد پارامترهای مدل‌های رگرسیون فازی از روش باز نمونه‌گیری بوت استرپ استفاده کردند. آن‌ها با باز نمونه‌گیری از خطاهای حاصل از یک برآورد اولیه که از طریق روش کمترین توان دوم خطا به دست آمده بود، نمونه‌های جدید را تولید کرده و برای هر کدام از آن نمونه‌ها برآورد کمترین توان دوم پارامترها را بدست آوردند. در واقع آنها ابتدا بر اساس مشاهدات اولیه، برآورد اولیه‌ای برای پارامترهای مدل به دست آورده و خطای برآورد را به صورت ۲.۲ محاسبه نمودند:

$$\hat{\epsilon}_i = \tilde{y}_i - \sum_{j=0}^n \hat{\beta}_j \tilde{x}_{ij}, \quad i = 1, \dots, n \quad (2.2)$$

که در آن $\hat{\beta}_j$ ها پارامترهای برآورد شده اولیه هستند و $\tilde{x}_{ij}$ مشاهدات ورودی و $\tilde{y}_i$ نیز پاسخ‌های مشاهده شده اولیه هستند.

^۱Efron

^۲Bootstrap

در این مرحله با روش با جایگذاری تعداد B نمونه بوت استرپی برای خطای برآورد به صورت $(\hat{\epsilon}^{(1)}, \dots, \hat{\epsilon}^{(B)})$ تولید و با اضافه کردن هر کدام از این خطاها به $\sum_{j=0}^n \hat{\beta}_j \tilde{x}_{ij}$ تعداد B مشاهده جدید را تولید کردند.

اکنون با این مشاهدات جدید، تعداد B مدل و B مجموعه پارامتر برآورد شده برای ضرایب مدل تولید می‌شود که هر کدام از آنها را با $\hat{\beta}^{(d)}$ $d = 1, 2, \dots, B$ نشان می‌دهیم. حال میانگین این B برآورد، برای هر پارامتر، به عنوان برآورد بوت استرپ به صورت زیر بدست خواهد آمد:

$$\hat{\beta}^B = \frac{1}{B} \min \sum_{d=1}^B \hat{\beta}^{(d)}.$$

آنها نشان دادند که مجموع خطاهای برآورد شده برای مدل تحت این روش در قیاس با روش‌های دیگر کمتر است. در سال‌های ۲۰۱۲، ۲۰۱۳ و ۲۰۱۵ نیز از روش بوت‌استرپ برای تهیه فواصل اطمینان مربوط به پارامترهای مدل رگرسیونی استفاده گردید. ساختار الگوریتمی برای اجرای روش بوت‌استرپ و استفاده از آن، با مراجعه به الگوریتم ۴ قابل مشاهده است.

در اینجا و با استفاده از این ساختار و روش باز نمونه‌گیری بوت‌استرپ و البته با استفاده از متر ارائه شده در [۸۱]، برای مجموعه داده‌های سوم ۳.۲، برآوردهای بوت استرپی مدل را به صورت زیر بدست آورده‌ایم.

$$\hat{\tilde{y}}_i = (3/575, \circ/\circ\circ)_T + (\circ/541, \circ/973)_T \tilde{x}_i$$

همچنین ارزیابی مدل تولید شده توسط معیارهای ارزیابی مدل که در بخش ۳.۱ تعریف شده‌اند در جدول ۵.۲ ارائه شده است.

جدول ۵.۲: معیارهای ارزیابی برای مدل بوت استرپی، در مجموعه داده‌های سوم ۳.۲.

مدل	SIM			MPE		
	S_3	S_2	S_1	P_3	P_2	P_1
بوت استرپ	۰/۶۱۲۹	۰/۶۱۲۳	۰/۶۱۲۷	۰/۷۳۶۷	۰/۷۳۶۹	۰/۷۳۶۸

۴.۲ باز نمونه‌گیری در رگرسیون فازی از طریق جک نایف بعد از بوت استرپ

خطای استاندارد و اریبی حاصل از روش بوت استرپ در بخش ۳.۲ بواسطه تصادفی بودن نمونه‌های تولیدی آن یک متغیر تصادفی خواهد بود. لذا ما در این بخش از رساله با استفاده از روش جک نایف بعد از بوت استرپ^۳ (JB) فازی علاوه بر رفع این مشکل نشان می‌دهیم که هم به لحاظ بهینگی مدل یعنی با استفاده از معیارهای نیکویی برازش و هم از دیدگاه استنباطی و با استفاده از فواصل اطمینان در جهت بررسی آزمون فرض ضرایب و همین طور نشان دادن خاصیت استوار بودن در مواجهه با نقاط دور افتاده عملکرد بهتری نسبت به روش بوت استرپ و دیگر روش‌های مدل‌بندی فازی دارد [۷۹]. در واقع ما اثر بخشی و عملکرد روش JB را از سه منظر بهینه کردن مدل، تولید برآوردی مناسب برای خطای استاندارد و خاصیت استوار بودن مدل تولیدی نشان می‌دهیم. فرض کنید مدل رگرسیون فازی به صورت زیر باشد:

$$\tilde{y}_i = \tilde{\mathbf{x}}_i^T \tilde{\boldsymbol{\beta}} + \tilde{\epsilon}_i, \quad i = 1, 2, \dots, n \quad (3.2)$$

که در آن $\tilde{y}_i$ i امین پاسخ مشاهده شده، $\tilde{\mathbf{x}}_i = (1, \tilde{x}_{i1}, \tilde{x}_{i2}, \dots, \tilde{x}_{ip})^T \in \mathbb{R}^{p+1}$ i مین بردار متغیرهای پیشگو، $\tilde{\boldsymbol{\beta}} = (\tilde{\beta}_0, \tilde{\beta}_1, \dots, \tilde{\beta}_p)^T$ ، بردار ضرایب مدل رگرسیونی فازی و ۱ در ماتریس طرح یک عدد فازی با پهنای صفر است. همچنین جمله خطای مدل به صورت $\tilde{\epsilon}_i$ آمده است. از آنجایی که برای اجرای روش بوت استرپ و به دنبال آن جک نایف بعد از بوت استرپ (JB)، با استفاده از ساختار باز نمونه‌گیری از خطاها، نیاز به یک برآورد اولیه برای بدست آوردن بردار خطا داریم، لذا ابتدا با استفاده از مدل و متر ارائه شده در [۸۱] به عنوان متر کاربردی ما، یک برآورد اولیه از مدل رگرسیونی می‌سازیم. به طور خلاصه روش (JB) یک روش باز نمونه‌گیری است که بر روی نمونه‌های بوت استرپی تولید شده اجرا می‌گردد. این روش بر اساس ساختار پیشنهادی در [۴۸] که در بخش چهارم از فصل نوزدهم کتاب آمده است اجرا می‌گردد. البته توجه داشته باشید که نمونه‌های بوت استرپ تولید شده در مدل رگرسیونی با استفاده از بوت استرپ خطاها حاصل شده‌اند. الگوریتم ۱ اجرای روش جک نایف بعد از بوت استرپ را نشان می‌دهد.

^۳ Jackknife-after-Bootstrap

تبصره ۱.۴.۲. برای هر نقطه $i, i = 1, 2, \dots, n$ از مشاهدات نمونه به میزان $(1 - \frac{1}{B})^B$ احتمال وجود دارد که مشاهده i ام در نمونه بوت استرپ به اندازه B ظاهر نشود [۴۸]. این احتمال به طور مجانبی وقتی $B \rightarrow \infty$ میل کند برابر با e^{-1} خواهد بود (e عدد نپر است که تقریباً برابر 2.718 می‌باشد). بنابراین در یک نمونه بوت استرپ به اندازه B ، در مرحله ۵ از الگوریتم ۱، $m \approx \frac{B}{e}$ تعداد ستون وجود دارد که در آن‌ها مشاهده i ام ظاهر نشده است. این ستون‌ها، برای اعمال روش JB بکار می‌رود.

الگوریتم ۱ الگوریتم محاسباتی روش JB

گام ۱: B را به عنوان تعداد نمونه بوت‌استرپی در نظر می‌گیریم.
گام ۲: تولید B نمونه تصادفی از خطای برآورد شده تحت یک برآورد اولیه به صورت زیر

$$\hat{\epsilon}^{(j)} = \tilde{\mathbf{Y}}^{(j)} - \hat{\mathbf{Y}}^{(j)}, \quad j = 1, \dots, B$$

که در آن j بیانگر شماره نمونه بوت‌استرپی است.
گام ۳: محاسبه پاسخ‌های مشاهده شده جدید با استفاده از نمونه‌های گام دوم و فرمول زیر

$$\hat{\mathbf{Y}}^{*(j)} = \hat{\epsilon}_i^{(j)} + \hat{\mathbf{Y}}_i^{(j)}, \quad j = 1, \dots, B$$

گام ۴: (برآورد بوت‌استرپ) برای هر کدام از B نمونه تولید شده، برآورد پارامترهای مدل رگرسیونی فازی را به‌دست می‌آوریم. اکنون برآورد بوت‌استرپی به صورت زیر خواهد بود:

$$\hat{\beta}^B = \frac{1}{B} \min_{j=1} \sum_{j=1}^B \tilde{\beta}^{(j)}.$$

که در آن، $\tilde{\beta}^{(j)}$ بیانگر برآورد مدل رگرسیونی مربوط به نمونه بوت‌استرپی j ام است.
گام ۵: (برآوردگر JB) با تقسیم $\{\hat{\mathbf{Y}}^{*(1)}, \dots, \hat{\mathbf{Y}}^{*(B)}\}$ به n گروه و هر کدام شامل $m \approx \frac{B}{e}$ ستون برای هر مشاهده از داده‌های نمونه پایه (به تبصره ۱.۴.۲ مراجعه شود). برای هر گروه، i امین سطر متناظر با آن مشاهده را در ماتریس طرح حذف می‌شود. سپس برآورد $\tilde{\beta}$ را برای هر کدام از m ستون به‌دست آورده و با $\hat{\beta}_{hi}$ نمایش می‌دهیم. اکنون برآورد JB عبارتست از :

$$\hat{\beta}^{JB} = \frac{1}{nm} \left(\sum_{i=1}^n \sum_{h=1}^m \hat{\beta}_{hi} \right)$$

۱.۴.۲ معیارهای ارزیابی مدل

اکنون پس از محاسبه برآوردهای JB و بوت استرپ، برای ارزیابی مدل‌های تولید شده از معیارهای MPE و SIM استفاده خواهیم کرد که در آ.۳، تشریح شده‌اند. همچنین در این تحقیق از فاصله اطمینان t -بوت استرپی برای استفاده در بخش استنباطی استفاده شده است که روش ساخت آن در الگوریتم ۲ تشریح می‌گردد.

الگوریتم ۲ فاصله اطمینان t -بوت استرپی (CI)

گام ۱: یک نمونه بوت استرپی به اندازه B در نظر بگیرید.

گام ۲: برآوردهای بوت استرپی $\hat{\beta}_k^{(j)} = (\hat{m}_{\hat{\beta}_k^{(j)}}, \hat{l}_{\hat{\beta}_k^{(j)}}, \hat{r}_{\hat{\beta}_k^{(j)}})_{LR}$ را با استفاده از رابطه

$$\tilde{z}_k = \left(\frac{\hat{m}_{\hat{\beta}_k^{(j)}} - \hat{m}_{\hat{\beta}_k^B}}{se(\hat{m}_{\hat{\beta}_k^{(j)}})}, \frac{\hat{l}_{\hat{\beta}_k^{(j)}} - \hat{l}_{\hat{\beta}_k^B}}{se(\hat{l}_{\hat{\beta}_k^{(j)}})}, \frac{\hat{r}_{\hat{\beta}_k^{(j)}} - \hat{r}_{\hat{\beta}_k^B}}{se(\hat{r}_{\hat{\beta}_k^{(j)}})} \right), k = 0, 1, \dots, p, j = 1, \dots, B$$

استاندارد کنید.

گام ۳: چندک‌های $\frac{\alpha}{2}$ و $(1 - \frac{\alpha}{2})$ را در مقادیر $\tilde{z}_k$ بدست آورده و به ترتیب با $\hat{t}_{n-1, \frac{\alpha}{2}}^{(*)}$ و $\hat{t}_{n-1, 1-\frac{\alpha}{2}}^{(*)}$ که در آن $* = m, l, r$ است مشخص کنید.

گام ۴: فواصل اطمینان را برای هر برآورد بوت استرپ به صورت زیر به دست آورید.

$$\begin{cases} \left[\hat{m}_{\hat{\beta}_k^{(j)}} + t_{n-1, \frac{\alpha}{2}}^{(m)} se(\hat{m}_{\hat{\beta}_k^{(j)}}), \hat{m}_{\hat{\beta}_k^{(j)}} + t_{n-1, 1-\frac{\alpha}{2}}^{(m)} se(\hat{m}_{\hat{\beta}_k^{(j)}}) \right] \\ \left[\hat{l}_{\hat{\beta}_k^{(j)}} + t_{n-1, \frac{\alpha}{2}}^{(l)} se(\hat{l}_{\hat{\beta}_k^{(j)}}), \hat{l}_{\hat{\beta}_k^{(j)}} + t_{n-1, 1-\frac{\alpha}{2}}^{(l)} se(\hat{l}_{\hat{\beta}_k^{(j)}}) \right] \\ \left[\hat{r}_{\hat{\beta}_k^{(j)}} + t_{n-1, \frac{\alpha}{2}}^{(r)} se(\hat{r}_{\hat{\beta}_k^{(j)}}), \hat{r}_{\hat{\beta}_k^{(j)}} + t_{n-1, 1-\frac{\alpha}{2}}^{(r)} se(\hat{r}_{\hat{\beta}_k^{(j)}}) \right] \end{cases}$$

توجه داشته باشید که ما در اینجا از ساختار فاصله اطمینان t -بوت استرپی استفاده کرده‌ایم که در [۱۳۴] پیشنهاد شده است.

حال پس از بدست آوردن برآورد JB برای پارامترهای مدل رگرسیونی فازی، در جهت بررسی و نمایش اثرگذاری این روش در بهبود مدل، رفع مشکل برآوردگر بوت استرپ و به علاوه تاثیر این روش در بررسی آزمون فرض پارامترهای برآورد شده مدل، پس از عنوان معیارهای مقایسه مربوطه مثال‌هایی واقعی را ارائه کرده و اثربخشی این روش را نشان خواهیم داد. همچنین برآوردانحراف استاندارد بوت استرپ هر پارامتر مانند $\hat{\beta}_k$ عبارتست از

$$se_{\hat{\beta}_k} = \left\{ \sqrt{\frac{\sum_{j=1}^B d^2(\hat{\beta}_j, \hat{\beta}_k^B)}{B-1}} \right\}^{\frac{1}{2}}, \quad k = 0, 1, \dots, p \quad (4.2)$$

همچنین اگر اعداد فازی $\hat{\beta}_k$ و $\hat{\beta}_k^B$ به صورت

زیر حاصل می‌گردد: $\hat{\beta}_k^B = (\hat{\beta}_k^B, s_{\hat{\beta}_k^B})_T$ و $\hat{\beta}_k = (\hat{\beta}_k, s_{\hat{\beta}_k})_T$ تعریف شوند، آنگاه d^2 از طریق فاصله اقلیدسی به صورت

$$d^2(\hat{\beta}_k, \hat{\beta}_k^B) = (\hat{\beta}_k - \hat{\beta}_k^B)^2 + (s_{\hat{\beta}_k} - s_{\hat{\beta}_k^B})^2$$

۲.۴.۲ مثال ساختگی ساکاوا و یانو (۱۹۹۲)

داده‌های موجود در این مثال که در بخش مثال‌های انگیزشی به عنوان مجموعه داده سوم در ۳.۲ آمده است، اولین بار توسط [۱۱۶] ارائه گردید و در ادامه توسط [۷۸، ۹۶، ۸۱] مورد استفاده و تجزیه و تحلیل قرار گرفت. ما سه مدل فوق را به ترتیب با عناوین “KT”، “KC” و “NN” در نظر گرفته و با اعمال روش JB بر روی آن‌ها به بررسی ارزیابی بهینگی مدل پیشنهادی خود پرداخته‌ایم. نتایج حاصله در جدول ۶.۲ آمده است.

جدول ۶.۲: معیارهای ارزیابی مدل قبل و بعد از بکارگیری روش JB برای مجموعه داده‌های سوم.

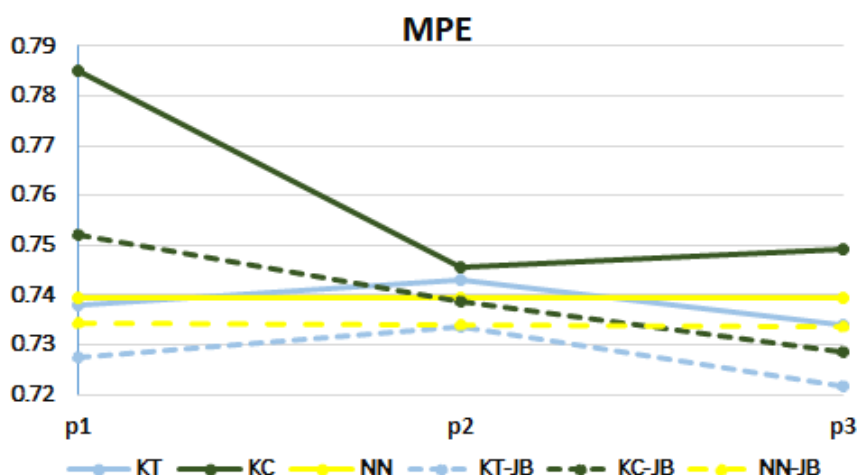
مدل قبل از اعمال روش JB			مدل پس از اعمال روش JB			معیار ارزیابی
NN	KC	KT	NN	KC	KT	
۰/۷۳۴۲	۰/۷۵۲۲	۰/۷۲۷۶	۰/۷۳۹۴	۰/۷۸۵۰	۰/۷۳۷۸	P_1
۰/۷۳۴۱	۰/۷۳۸۸	۰/۷۳۳۶	۰/۷۳۹۵	۰/۷۴۵۴	۰/۷۴۳۱	P_2
۰/۷۳۳۶	۰/۷۲۸۶	۰/۷۲۱۵	۰/۷۳۹۳	۰/۷۴۹۱	۰/۷۳۴۰	P_3
۰/۶۱۳۶	۰/۵۹۶۳	۰/۶۲۳۴	۰/۶۱۲۱	۰/۵۷۱۴	۰/۶۱۱۸	S_1
۰/۶۱۳۸	۰/۶۱۲۰	۰/۶۲۲۱	۰/۶۱۱۵	۰/۵۹۱۵	۰/۶۱۰۴	S_2
۰/۶۱۴۲	۰/۶۱۲۷	۰/۶۲۸۶	۰/۶۱۲۵	۰/۵۸۷۱	۰/۶۱۳۲	S_3

با توجه نتایج جدول ۶.۲، با وجود کم بودن مقادیر MPE در روش‌های سه گانه آزموده شده، این شاخص‌ها در مدل‌های متناظر با آن‌ها که روش JB در آن‌ها اعمال شده است کاهش داشته است. همچنین مقادیر شباهت نیز در مقابل مدل‌های مبنا افزایش یافته است. این برتری در نمودار ۱.۲ نمایش داده شده است. همچنین ضرایب مدل‌های سه گانه و مدل‌های JB متناظر آن‌ها در جدول زیر ارائه شده است. توجه داشته باشید که در مدل “KC” آخرین جمله اضافه شده، جمله خطا است (برای جزئیات بیشتر به [۷۸] مراجعه کنید).

۴۶ برآورد ضرایب مدل رگرسیون فازی به روش جک نایف بعد از بوت استرپ

جدول ۷.۲: ضرایب برآورد شده برای مدل های مبنا و مدل های JB متناظر آنها در مجموعه داده های سوم.

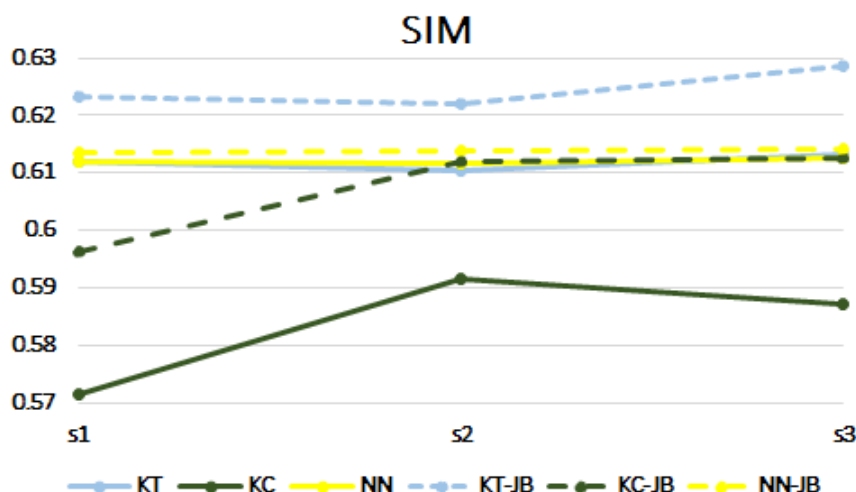
مدل ها	ضرایب مدل		
	$\hat{\beta}_0$	$\hat{\beta}_1$	$\hat{\epsilon}_i$
KT	$(4/103, 0/500)_T$	$(0/444, 0/080)_T$	---
$KT - JB$	$(4/071, 0/328)_T$	$(0/441, 0/022)_T$	---
KC	$(3/565, 0/000)_T$	$(0/522, 0/000)_T$	$(-0/011, 0/945)_T$
$KC - JB$	$(3/557, 0/000)_T$	$(0/518, 0/000)_T$	$(-0/013, 0/723)_T$
NN	$(3/577, 0/000)_T$	$(0/547, 1/000)_T$	---
$NN - JB$	$(3/568, 0/000)_T$	$(0/535, 0/968)_T$	---



(آ) مقدار MPE

همانطور که در نمودار ۱.۲ مشاهده می شود، اثر روش JB بر روی مدل های سه گانه باعث کاهش مقادیر MPE و افزایش SIM در تقابل با مدل های متناظر پایه شده است. در واقع خطوط نقطه چین در رنگ های مختلف بیانگر عملکرد مدلها پس از اعمال روش JB است.

اکنون در ادامه ارزیابی بهینگی روش پیشنهادی، فواصل اطمینان بوت استرپی به صورت زیر ارائه می گردد. این اقدام از دو منظر بکار رفته است: در ابتدا مقایسه طول بازه تولید شده توسط دو روش بوت استرپ و JB را انجام می دهیم و سپس بکارگیری این بازه ها را در بررسی آزمون فرض ضرایب دو مدل خواهیم داشت.



(ب) مقدار SIM

شکل ۱.۲: مقایسه شهودی مدل‌های آزموده با مدل‌های JB متناظر آن‌ها برای مجموعه داده‌های سوم. (آ) مقدار MPE (ب) مقدار SIM.

جدول ۸.۲: فواصل اطمینان بوت استرپی برای پارامترها.

روش	مولفه	$\tilde{\beta}_0$	$\tilde{\beta}_1$
بوت استرپ	مرکز	$[۲/۱۶, ۵/۴۲]$	$[۰/۲۲۸, ۰/۶۵۵]$
		$(۳/۲۶۸)$	$(۰/۴۲۸)$
	پهنا	$[۰/۰۰۰, ۰/۶۵۶]$	$[۰/۰۰۰, ۰/۰۴۵]$
		$(۰/۶۵۶)$	$(۰/۰۴۵)$
JB	مرکز	$[۳/۸۹۶, ۴/۲۴۶]$	$[۰/۴۰۳, ۰/۴۶۸]$
		$(۰/۳۴۹)$	$(۰/۰۶۵)$
	پهنا	$[۰/۰۸۴, ۰/۳۷۴]$	$[۰/۰۱۲, ۰/۰۲۶]$
		$(۰/۲۹۰)$	$(۰/۰۱۴)$

جدول ۸.۲ بیانگر فواصل اطمینان ۹۵ درصدی است که برای دو روش بوت استرپ و JB تولید شده است. با توجه به نتایج این جدول، طول فواصل اطمینان بدست آمده که در پرانتز نشان داده شده است برای تمام ضرایب، در مراکز و پهناها، در روش JB بسیار کوتاه‌تر از فواصل متناظر آن‌ها در روش بوت استرپ است که این بواسطه کاهش خطای استاندارد در روش JB است.

نکته قابل تامل در این نتایج آن است که بنابر خاصیت مثبت بودن پهناها در صورتی که حد پایین بازه اطمینان مربوط به پهناهای یک ضریب منفی شود، به ناچار آن را صفر در نظر

می‌گیریم که این حرکت باعث کاهش طول بازه اطمینان می‌شود. این مسئله در مورد پهنای ضرایب برآورد شده در مدل بوت استرپ رخ داده است حال آنکه در روش JB با این مسئله روبرو نشده‌ایم. با این حال بازهم طول بازه اطمینان تولید شده در روش JB کوتاه‌تر از طول بازه اطمینان در روش بوت استرپ است. در نمودار شماره ۲.۲ به صورت شهودی می‌توان نتایج بدست آمده در جدول ۸.۲ را مشاهده نمود.

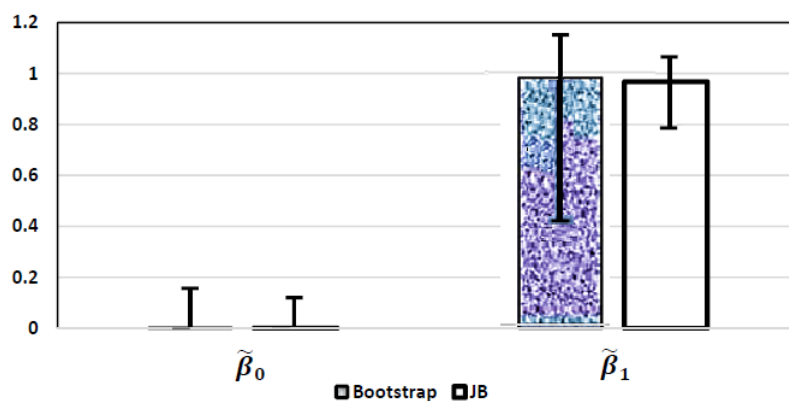
اکنون با استفاده از فواصل اطمینان به‌دست آمده به بررسی آزمون فرضیه ضرایب که به صورت زیر ارائه شده است می‌پردازیم:

$$H_0 : \tilde{\beta}_i \approx \{0\}, \quad i = 0, 1,$$

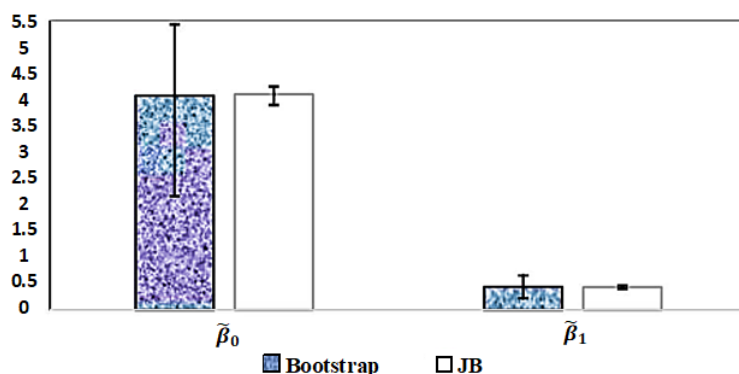
که در آن $\{0\}$ عدد فازی صفر با پهنای صفر و $\approx$ به مفهوم تقریباً برابر بکار رفته است. برای مطالعه این آزمون فرضیه، مرکز و پهنای ضرایب را به طور جداگانه در نظر می‌گیریم. یعنی، بعد از ایجاد فواصل اطمینان‌های مربوطه دقیقاً همانطور که در جدول ۸.۲ ارائه شده است، طبق استنباط‌های حاصل از فاصله اطمینان، اگر این بازه‌ها شامل مقدار صفر باشند، فرضیه صفر را می‌پذیریم.

از سوی دیگر، آنجا که برآورد ضرایب فازی برای برخی از متغیرهای ورودی شامل حالت‌هایی چون $(0,0)_T$ یا $(0,s)_T$ باشد، نتیجه می‌گیریم که اگر، حداقل برای مراکز، بازه‌های مرتبط با آن‌ها شامل صفر باشند آن‌ها را ”حدوداً صفر” در نظر گرفته و از مدل برازش شده حذف می‌کنیم زیرا آن‌ها متغیرهای ناموثر هستند (برای جزئیات بیشتر به [۶۵] مراجعه کنید).

در اینجا، برای پارامترهای برآورد شده براساس فواصل اطمینان، در سطح معنی داری ۵ درصد، ضرایب بدست آمده در هر دو مدل پذیرفته می‌شوند.



(آ) پهنای ضرایب



(ب) مرکز ضرایب

شکل ۲.۲: مقایسه شهودی فواصل اطمینان بین بوت استرپ و JB (آ) پهنای ضرایب (ب) مرکز ضرایب.

۳.۴.۲ مثال زمان پاسخگویی

این داده‌ها که در بخش مثال‌های انگیزشی به عنوان داده‌های اول که در ۱.۲ آمده است، دارای سه متغیر ورودی غیر فازی و یک متغیر فازی مثلثی متقارن است. در اینجا، ما از مدل‌های [۸۱]، [۱۵۰]، [۳۶] و [۹۳] برای مقایسه با مدل پیشنهادی خود که بر اساس متر D_3 که در بخش ۱.۲.۱ ضمیمه تشریح شده‌اند، تحت روش JB تولید می‌شود، استفاده کرده‌ایم و آن‌ها را به ترتیب با «مدل KT»، «مدل ZFL»، «مدل CDGS» و «مدل MT» علامت‌گذاری کرده‌ایم.

در این مسیر علاوه بر ارائه معیارهای نیکویی برازش (GOF)^۴ در خصوص ارزیابی بهینگی مدل پیشنهادی در مقابل دیگر مدل‌های مورد آزمایش و تولید فواصل اطمینان و بررسی آزمون فرض ضرایب، خاصیت استوار بودن روش JB را هم در مواجهه با نقاط دور افتاده مورد بررسی قرار داده‌ایم که در ادامه تشریح خواهد شد.

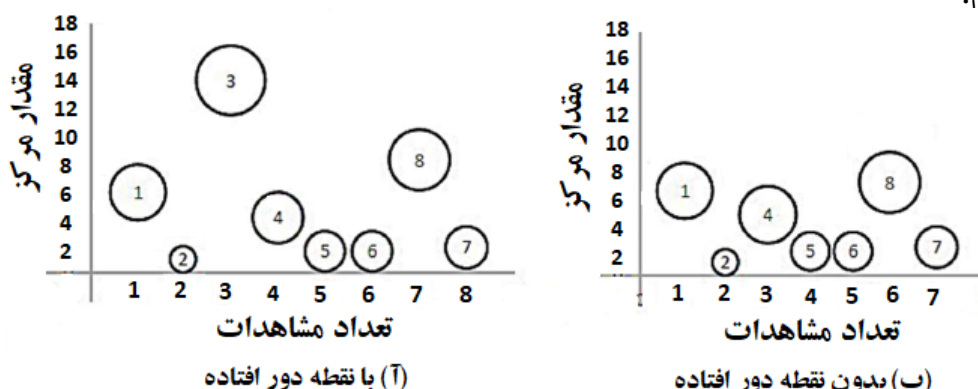
در یک سری از مجموعه مشاهدات که گاه مشاهداتی وجود دارند که به لحاظ شهودی در نمودارهای پراکنش موسوم به نمودار حبابی دورتر از دیگر مشاهدات قرار می‌گیرند و به لحاظ تکنیکی و بر اساس شاخص‌هایی چون فاصله کوک^۵ (CD) [۳۵] و [۷۶] و از این دست، دارای فاصله دورتر نسبت به دیگر مشاهدات را نشان می‌دهند که به عنوان نقاط دور افتاده (پرت) در نظر گرفته می‌شوند. در این مثال هم علاوه بر نمودار پراکنش حبابی (نمودار ۳.۲)، نمودار فاصله کوک (نمودار ۴.۲) نیز ارائه شده است.

تبصره ۲.۴.۲. از آنجایی که داده‌های مورد استفاده در متغیر پاسخ، فازی هستند برای بکارگیری روش تشخیصی CD آن‌ها را از طریق روش عمومی مرکز ثقل به داده‌های غیر فازی تبدیل

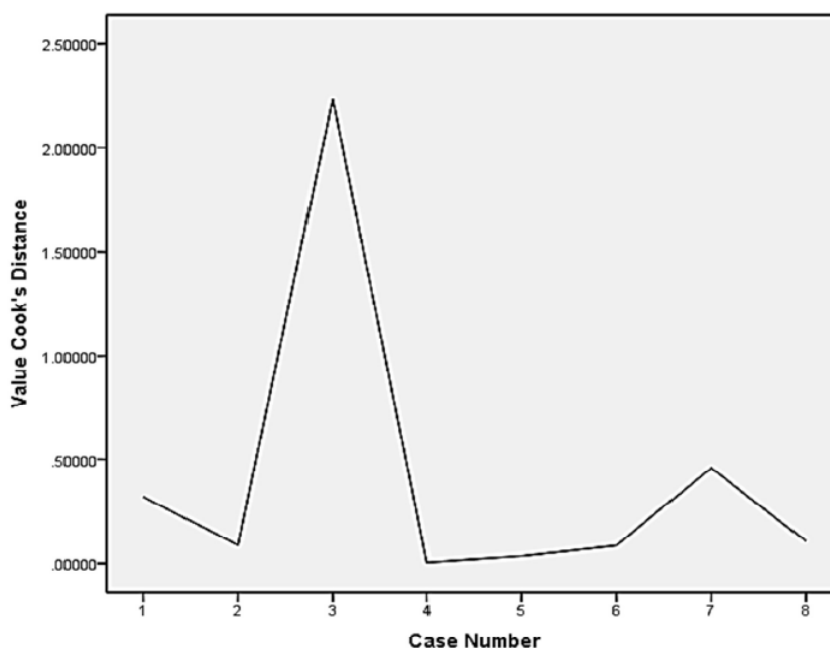
^۴ Goodness of fit

^۵ Cook distance

کرده‌ایم.



شکل ۳.۲: نمودار حبابی با و بدون مشاهده سوم به عنوان نقطه دور افتاده (آ) با نقطه دور افتاده (ب) بدون نقطه دور افتاده.



شکل ۴.۲: نمودار فاصله کوک در تشخیص نقطه دور افتاده.

با توجه به نمودار پراکنش حبابی مشاهده سوم را می‌توان به عنوان نقطه دور افتاده در این مجموعه داده در نظر گرفت. علاوه بر آن، از آنجایی که طبق نمودار ۴.۲ (که بر اساس مقادیر CD تولید شده‌اند) فاصله مشاهده سوم از دیگر مشاهدات بیشتر است و بخاطر شاخص ارزیابی که در [۱۱] پیشنهاد شده است، چون مقدار CD برای مشاهده سوم بیش از $\frac{4}{n}$ است به لحاظ تکنیکی نیز می‌توان این مشاهده را به عنوان نقطه دور افتاده در نظر گرفت. اکنون برای رسیدن به اهداف مورد نظر در این مثال یکبار مدل‌ها را با کل مشاهدات و یکبار بدون مشاهده دور افتاده برازش داده و ارزیابی‌ها و استنباط‌های مورد هدف را در جداول مربوطه ارائه خواهیم داد.

باز نمونه‌گیری در رگرسیون فازی از طریق جک نایف بعد از بوت استرپ ۵۱

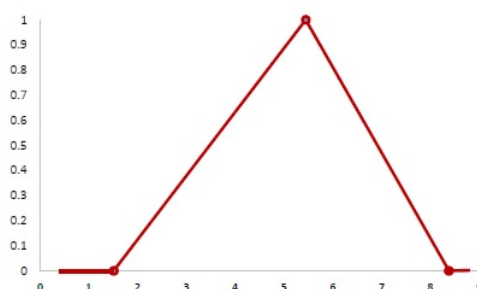
جدول ۹.۲: نتایج معیارهای ارزیابی GOF قبل و بعد از حذف نقطه دورافتاده.

مدل	شامل نقطه پرت						بدون نقطه پرت					
	S_3	S_2	S_1	P_3	P_2	P_1	S_3	S_2	S_1	P_3	P_2	P_1
KT	۰/۱۸۱۲	۰/۱۶۲۸	۰/۱۴۹۷	۴/۵۲۰۳	۵/۱۴۴۴	۵/۶۸۰۵	۰/۲۲۳۵	۰/۲۱۲۰	۰/۱۹۷۱	۳/۴۷۴۶	۳/۷۱۷۸	۴/۰۷۴۵
ZFL	۰/۱۸۴۳	۰/۱۶۸۰	۰/۱۴۹۸	۴/۴۲۵۷	۴/۹۵۱۱	۵/۶۷۹۹	۰/۲۳۶۸	۰/۲۱۴۰	۰/۲۰۹۱	۳/۲۲۲۵	۳/۶۷۲۶	۳/۷۸۳۱
$CDGS$	۰/۱۴۸۹	۰/۱۰۴۹	۰/۱۲۵۵	۵/۷۱۳۹	۸/۵۳۱۸	۶/۹۷۰۸	۰/۱۶۸۲	۰/۱۱۷۸	۰/۱۶۰۸	۴/۹۴۳۸	۷/۴۸۶۸	۵/۲۲۰۱
MT	۰/۱۵۷۵	۰/۱۴۶۸	۰/۱۴۸۸	۵/۳۴۸۱	۵/۸۱۱۹	۵/۷۱۸۴	۰/۲۱۸۱	۰/۲۰۷۹	۰/۱۷۹۲	۳/۵۸۵۲	۳/۸۲۶۴	۴/۵۵۸۱
$KT - JB$	۰/۳۱۵۱	۰/۲۸۹۷	۰/۳۵۳۰	۳/۱۱۰۱	۳/۵۸۵۵	۳/۲۳۶۶	۰/۳۲۳۱	۰/۳۰۱۱	۰/۳۵۷۳	۳/۰۰۱۵	۳/۳۸۷۲	۳/۱۴۹۳

بر اساس نتایج جدول ۹.۲، در بخش اول در بررسی با کل مشاهدات، مقادیر به‌دست آمده برای تمامی شاخص‌های ارزیابی MPE در روش JB کمتر از دیگر روش‌های مورد آزمایش است. این موضوع در مورد شاخص‌های شباهت SIM به گونه‌ای است که در آن‌ها میزان شباهت مدل JB نسبت به دیگر مدل‌ها بیشتر است. برای مثال در معیار P_1 مقدار عددی نتیجه شده برای مدل تولیدی تحت روش JB برابر با ۳/۲۳۶۶ بدست آمده است که این معیار در مدل‌های رقیب در بهترین حالت، مقدار عددی ۵/۶۷۹۹ است. در مورد شاخص شباهت نیز برای مثال، مقدار عددی بدست آمده در S_1 برای پاسخهای تولیدی به روش JB با پاسخهای مشاهده شده شباهتی به میزان ۳۵ درصد را نشان می‌دهد که این در مورد مدل‌های رقیب در بهترین حالت برابر با ۱۵ درصد است. بدین ترتیب نتایج حاصله بیانگر عملکرد بهتر مدل تولید شده توسط روش JB نسبت به مدل‌های رقیب است.

در بخش دوم یعنی مدل‌های تولید شده بدون وجود نقطه دور افتاده در نگاه اول، به وضوح کاهش مقادیر MPE و افزایش مقادیر SIM را می‌توان در چهار مدل رقیب مشاهده کرد با این تفاوت که مدل JB تولید شده در مرحله اول در هر دو شاخص ارزیابی، کماکان عملکرد بهتری نسبت به مدل‌های رقیب در حالت جدید دارد. همچنین مدل JB تولید شده در بخش دوم با اینکه نقطه دور افتاده را در اختیار ندارد و این امر باعث کاهش شاخص‌های ارزیابی MPE و افزایش مقادیر SIM شده است، اما این تغییرات نسبت به مقادیر به‌دست آمده در بخش اول بسیار کم است این در حالی است که در مدل‌های رقیب تغییرات ایجاد شده معنی دار است. این مطلب بیانگر وجود خاصیت استواری در روش JB و درواقع تاثیرپذیری ناچیز این روش از نقطه دورافتاده است.

بررسی شهودی مطالب استنباط شده از جدول فوق را می‌توان در نمودار ۶.۲ مشاهده کرد. همانطور که در نمودار ۶.۲ مشاهده می‌شود، کاهش شاخص MPE و افزایش شاخص SIM در روش پیشنهادی نسبت به مدل‌های رقیب در هر دو وضعیت حضور و عدم حضور نقطه دورافتاده در مشاهدات بیانگر عملکرد بهتر مدل JB در مقابل دیگر مدل‌های مورد مطالعه است. علاوه بر آن مدل‌های برازش شده در حالت عدم وجود نقطه دور افتاده با وجود تغییر محسوس و بهینه شدن همچنان در رقابت با مدل JB تولیدی با وجود نقطه پرت، مغلوب



شکل ۵.۲: تصویر پاسخ برآورد شده

شده‌اند. همچنین در مقایسه مدل JB در دو حالت، معیارهای بهینگی در وضعیت بدون حضور نقطه دورافتاده، چندان تغییر محسوسی را بواسطه استوار بودن این روش نسبت به نقاط پرت نشان نمی‌دهند.

مدل پیشنهادی در حضور نقطه دور افتاده عبارتست از:

$$\hat{y}_i = (1/8 \cdot 23, 0/1627)_T + (-0/3718, 0/9787)_{Tx_1} \\ + (-1/1449, 0/3227)_{Tx_2} + (0/4225, 0/1091)_{Tx_3}$$

بنابراین توجه به معادله رگرسیونی فوق که براساس مدل تولید شده تحت روش JB تولید شده است، اگر مدت تجربه خدمه در داخل اتاق کنترل ($x_1 = 1$) سال، مدت تجربه خدمه در خارج اتاق کنترل $x_2 = 1/5$ سال و مدت تحصیلات خدمه را $x_3 = 12$ سال در نظر بگیریم، آنگاه پاسخ برآورد شده توسط روش JB برابر است با $(5/4432, 2/9346)_T$ ، که یعنی با فرض اینکه مدت تجربه او در داخل اتاق کنترل یک سال و در خارج اتاق یک سال و نیم و مدت تحصیلات این فرد ۱۲ سال باشد، مدت زمان پاسخ شناختی خدمه داخل اتاق کنترل به یک رویداد غیر عادی یک عدد فازی مثلثی حدوداً $5/4432$ سال است.

نمودار؟؟ تصویر شهودی این عدد فازی است. اکنون برای نشان دادن اثر کاهشی روش JB در خطای استاندارد ضرایب برآورد شده و بررسی آزمون فرض آن‌ها، فواصل اطمینان بوت استرپی را برای مراکز و پهنای ضرایب در جدول ۱۰.۲ بدست می‌آوریم تا علاوه بر اثر بخشی و بهینگی، روش JB را در مقابل بوت استرپ و مقدمات بکارگیری در آزمون فرضیه زیر فراهم آورده باشیم.

با توجه به نتایج جدول ۱۰.۲، مشاهده می‌شود که طول فواصل اطمینان ۹۵ درصد تولید شده برای ضرایب و پهنای روش JB بسیار کوتاهتر از طول آن‌ها در مدل بوت استرپ است که این نه تنها بیانگر اثر بخشی روش JB در کاهش خطای استاندارد مربوط به ضرایب برآورد شده است بلکه به عنوان برآورد مناسبی برای خطای استاندارد ضرایب بوت استرپی، از افزایش بیش از اندازه آن جلوگیری می‌کند.

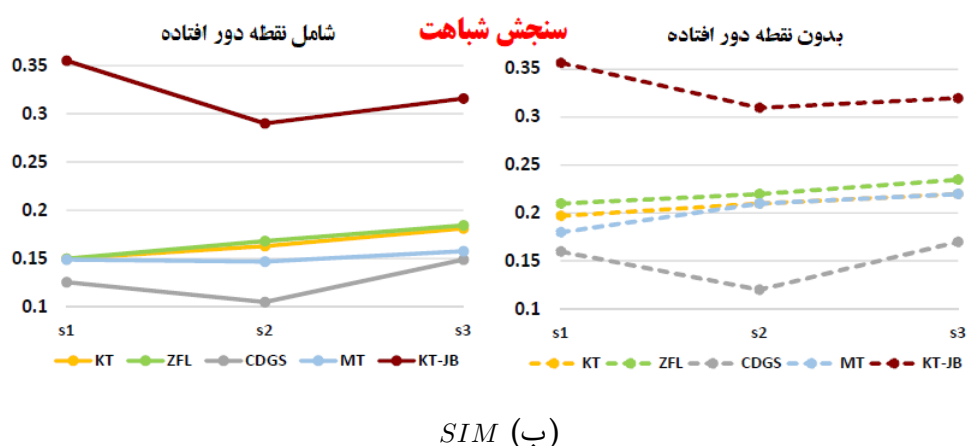
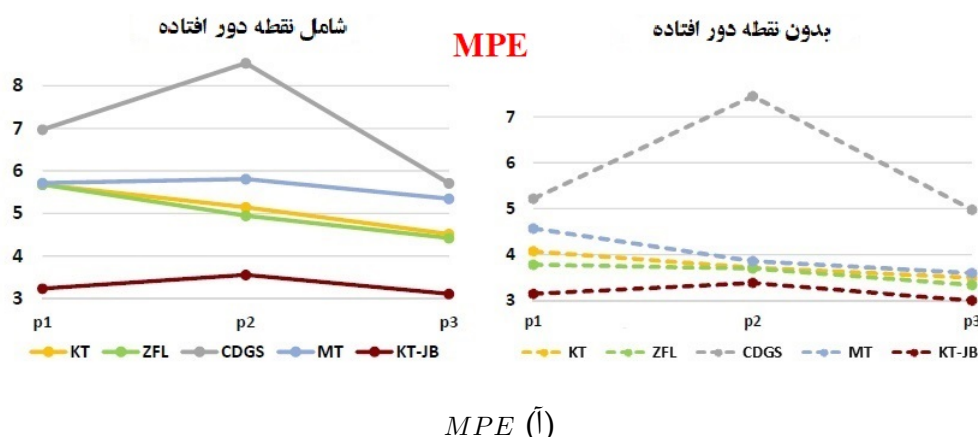
نکته دیگری که قابل توجه خواهد بود آن است که، بواسطه خاصیت نامنفی بودن پهنای، در صورتی که حد پایین فاصله اطمینان منفی شود به اجبار آن را صفر در نظر می‌گیریم که

باز نمونه‌گیری در رگرسیون فازی از طریق جک نایف بعد از بوت استرپ ۵۳

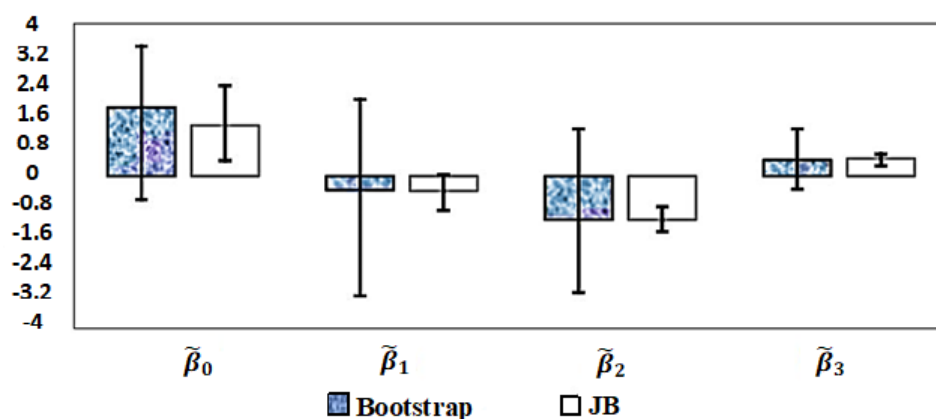
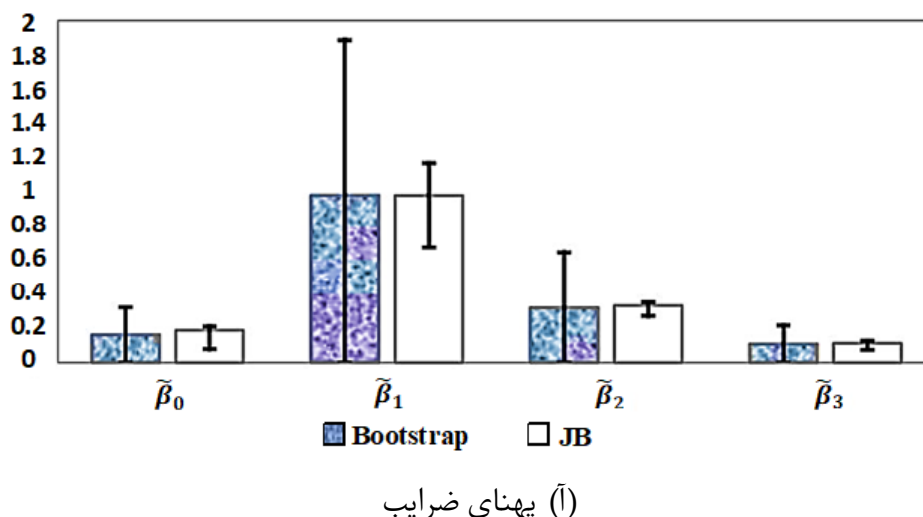
این امر در کاهش طول فاصله اطمینان پهنای ضرایب به نفع روش بوت استرپ است. اما در مثال فوق و طبق نتایج جدول ۱۰.۲ نه تنها روش JB این شرایط را دارا نیست بلکه با وجود این امتیاز به نفع بوت استرپ، باز هم طول بازه اطمینان پهنای در روش JB بسیار کوتاهتر از طول آن‌ها در روش بوت استرپ است. تصویر این برتری در نمودار ۷.۲ به وضوح قابل درک و مشاهده است.

جدول ۱۰.۲: فواصل اطمینان t -بوت استرپی. (مقادیر داخل پرانتز طول فاصله اطمینان است)

روش	مولفه	$\tilde{\beta}_0$	$\tilde{\beta}_1$	$\tilde{\beta}_2$	$\tilde{\beta}_3$
بوت استرپ	مرکز	$[-0.617, 3.406]$ (۴/۰۲۳)	$[-3.145, 2.012]$ (۵/۱۵۷)	$[-3.054, 1.238]$ (۴/۲۹۲)	$[-0.346, 0.921]$ (۱/۲۶۷)
	پهنا	$[0.000, 0.325]$ (۰/۳۲۵)	$[0.000, 1.886]$ (۱/۸۸۶)	$[0.000, 0.645]$ (۰/۶۴۵)	$[0.000, 0.218]$ (۰/۲۱۸)
	مرکز	$[0.403, 2.376]$ (۱/۹۷۳)	$[-0.902, 0.031]$ (۰/۹۳۳)	$[-1.259, -0.607]$ (۰/۶۵۲)	$[0.267, 0.577]$ (۰/۳۱۰)
	پهنا	$[0.078, 0.212]$ (۰/۱۳۴)	$[0.674, 1.166]$ (۰/۴۹۲)	$[0.273, 0.356]$ (۰/۰۸۳)	$[0.075, 0.127]$ (۰/۰۵۲)
JB	مرکز	$[0.403, 2.376]$ (۱/۹۷۳)	$[-0.902, 0.031]$ (۰/۹۳۳)	$[-1.259, -0.607]$ (۰/۶۵۲)	$[0.267, 0.577]$ (۰/۳۱۰)
	پهنا	$[0.078, 0.212]$ (۰/۱۳۴)	$[0.674, 1.166]$ (۰/۴۹۲)	$[0.273, 0.356]$ (۰/۰۸۳)	$[0.075, 0.127]$ (۰/۰۵۲)
	مرکز	$[0.403, 2.376]$ (۱/۹۷۳)	$[-0.902, 0.031]$ (۰/۹۳۳)	$[-1.259, -0.607]$ (۰/۶۵۲)	$[0.267, 0.577]$ (۰/۳۱۰)
	پهنا	$[0.078, 0.212]$ (۰/۱۳۴)	$[0.674, 1.166]$ (۰/۴۹۲)	$[0.273, 0.356]$ (۰/۰۸۳)	$[0.075, 0.127]$ (۰/۰۵۲)



شکل ۶.۲: تصویر عملکرد روش JB در مواجهه با نقطه دورافتاده: $(\bar{I})$ MPE (ب) SIM.



شکل ۷.۲: مقایسه فواصل اطمینان بدست آمده در دو روش بوت استرپ و JB. (آ) پهنای ضرایب (ب) مرکز ضرایب

اکنون پس از ارائه ضرایب مدل و تولید فواصل اطمینان بوت استرپی برای ضرایب هر دو روش به بررسی آزمون زیر می پردازیم.

$$H_0: \tilde{\beta}_i \approx \{0\}, \quad i = 0, 1, 2, 3$$

که در آن $\{0\}$ عدد فازی صفر با پهنای صفر و $\approx$ به مفهوم تقریباً برابر بکار رفته است. برای مطالعه این آزمون فرضیه، مرکز و گسترش ضرایب را به طور جداگانه در نظر می گیریم. یعنی، بعد از ایجاد فواصل اطمینانهای مربوطه دقیقاً همانطور که در جدول ۱۰.۲ ارائه شده است، طبق استنباطهای حاصل از فاصله اطمینان، اگر این بازه ها شامل مقدار صفر باشند، فرضیه صفر را می پذیریم. از سوی دیگر، آنجا که برآورد ضرایب فازی برای برخی از متغیرهای ورودی شامل حالت هایی چون $(0, 0)_T$ یا $(0, s)_T$ باشد، نتیجه می گیریم که اگر، حداقل برای مراکز، بازه های مرتبط با آنها شامل صفر باشند آنها را "حدوداً صفر" در نظر گرفته و از مدل

برازش شده حذف می‌کنیم زیرا آن‌ها متغیرهای ناموثر هستند (برای جزئیات بیشتر به [۶۵] مراجعه کنید).

بر این اساس و طبق نتایج جدول ۱۰.۲ و نمودار متناظر با این فواصل در ۷.۲، فرض صفر برای تمامی ضرایب مدل بوت استرپ در سطح معنی داری ۵ درصد پذیرفته شده و تمامی ضرایب از مدل حذف خواهند شد. اما در مورد نتایج بدست آمده برای مدل JB تنها متغیر x_1 است که بازه اطمینان ۹۵ درصدی شامل صفر دارد که از مدل خارج می‌شود.

۴.۴.۲ نتیجه‌گیری

در مواجهه با نمونه‌های کوچک روش بوت استرپ از کارایی بالایی در تخمین ضرایب مدل رگرسیونی فازی برخوردار است. با این حال خطای استاندارد برآورد، تصادفی می‌شود. برای مقابله با این نقص و بهبود تخمین، در این مقاله، ما روش JB را پیشنهاد کردیم. در این فصل با بکارگیری روش JB نتایج زیر حاصل گردید.

(الف) استفاده از روش JB به عنوان یک روش ساده و کارآمد در برآورد انحراف معیار پارامترهای مدل بواسطه مشکل تصادفی بودن انحراف معیار تولید شده در روش بوت استرپ که در [۴۸] به تفصیل به آن اشاره و روش JB برای رفع آن پیشنهاد شده است.

(ب) برازش بهتر مدل‌های رگرسیونی فازی در مقابل دیگر مدل‌های آزمایش شده که در بخش ۲.۴.۲ و ۳.۴.۲ نشان داده شده است.

(ج) کوتاه شدن فواصل اطمینان بوت استرپی در روش JB در مقابل روش بوت استرپ به‌واسطه کاهش خطای استاندارد ضرایب برآوردشده.

(د) دارا بودن خاصیت استواری در مواجهه با اثر نقاط پرت با حفظ برتری از لحاظ بهینگی در مقابل مدل‌های دیگر که در بخش ۳.۴.۲ مورد بررسی قرار گرفته است.

لازم به ذکر است که روش JB فقط در مواردی قابل اعتماد است که تعداد نمونه بوت استرپ به اندازه کافی بزرگ یعنی بیش از ۲۰۰ تکرار باشد و برای تکرارهای کمتر از ۵۰ دارای بیش برآورد است. در حالت غیر هموار نیز با استفاده از روش جک نایف d - حذفی قابل اجرا است. (به بخش ۱۹.۴، صفحه ۲۸۰-۲۷۵ از [۴۸] مراجعه کنید). مطالعات آینده می‌تواند گسترش JB را در مدل‌سازی غیر هموار با استفاده از روش جک نایف d - حذفی در نظر گرفته شود. روش ما همچنین می‌تواند برای مجموعه‌های فازی - بازه‌ای اعمال شود.

فصل ۳

مدل‌های رگرسیونی انقباضی فازی

۱.۳ مقدمه

همانطور که در ابتدای رساله بیان گردید، در مدل‌بندی رگرسیونی کار بر روی برآوردهای یک دیدگاه جهت افزایش بهینگی مدل محسوب می‌شوند. این کار از دو منظر قابل بررسی است. در این بخش به بررسی و اثر بهبود دهندگی روش‌های انقباضی به عنوان اولین رویکرد ذکر شده در این خصوص می‌پردازیم. در این بحث، یک برآوردهگر اریب ولی با واریانس کمتر از برآوردهگر نااریب تولید می‌شود. در واقع این برآوردهگر در مقایسه برآوردهگر نااریب در توازن بین میزان افزایش (کاهش) اریبی و واریانس، میانگین توان دوم خطای کمتری را تولید می‌کند. توضیحات جامع‌تر در نحوه عملکرد برآوردهگر انقباضی در پیوست ”ب“ بخش ۱.۱ و ۲.۱ آورده شده است.

روش انقباضی روشی است که با ایجاد تدریجی اریبی در برآورد، به طور صریح یا ضمنی، میانگین خطای پیش‌بینی^۱ (MPE) را بهبود می‌بخشد. از طرف دیگر، در مدل‌سازی داده‌های پیچیده، روش‌های بهبود یافته می‌توانند مدل‌های سازگارتر را از نظر عملکرد تولید کنند. اثر بخشی برخی از روش‌های بهبود یافته علیرغم ماهیت غیر فازی بودن آن‌ها را می‌توان در [۱۵۲]، [۳۷] و [۳۸] مشاهده کرد. در اینجا، ما از استراتژی انقباض به عنوان یک روش بهبود یافته در مدل‌سازی فازی استفاده می‌کنیم. این استراتژی برآورد خام (اولیه) را که ممکن است از

^۱Mean prediction error

روش‌های مختلفی مانند ماکزیمم درست‌نمایی، کمترین توان دوم خطا و یا رویکردهای نوآورانه بدست آید، به سمت یک برآورد هدف منقبض می‌کند و در نهایت یک برآورد ترکیبی با MPE کوچکتر در مقایسه با برآورد خام ارائه می‌دهد.

در این بخش، ما برای تخمین ضرایب مدل رگرسیون فازی از روش برآورد انقباضی معروف به برآورد انقباضی نوع استاین [۱۲۳] بهره برده‌ایم. روش‌های انقباضی اغلب هنگامی که اطلاعات غیر نمونه‌ای برای بهبود پیش‌بینی مدل در دسترس است، استفاده می‌شود. به عبارت دیگر، برای هر برآوردگر مانند X ، برآوردگر انقباضی به صورت $X + \omega g(X)$ است که در آن ω را ثابت انقباضی گوئیم و $0 < \omega < 1$ است. در واقع برای بهبود برآورد بدون تغییر در شرایط یا توابع هدف و همچنین با هزینه محاسباتی و زمانی بسیار ناچیز می‌توان برآوردگر تولید شده X را به سمت مقدار عددی خاصی که ممکن است از سوی اطلاعات غیر نمونه‌ای آمده باشد منقبض کرد. این انقباض در برآورد انقباضی نوع استاین به سمت عدد ضفر انجام می‌پذیرد.

در ادامه موارد زیر قابل مطالعه است:

الف) استفاده از روش انقباضی نوع استاین در مدل رگرسیونی فازی

ب) ارائه روشی بنام روش انقباضی فازی یکپارچه^۲ (IFS) در برآورد پارامترها و تولید یک مدل رگرسیونی فازی که در مقایسه با روش‌ها و مدل‌های دیگر تحت معیارهای ارزیابی‌های مختلف، دارای عملکرد بهتری است بخش ۳.۳.

ج) بررسی آزمون فرض ضرایب بر اساس یک آزمون فرض فازی و تحت روش بوت استرپ.

۱.۱.۳ مدل انقباضی نوع استاین

در این قسمت با توجه به مطالب ذکر شده، مدل انقباضی نوع استاین فازی را ارائه می‌کنیم. فرض کنید مدل رگرسیونی فازی به صورت ۱.۳ باشد:

$$\tilde{y}_i = \oplus_{j=0}^p (\tilde{\beta}_j \otimes \tilde{x}_{ij}) \oplus \tilde{\epsilon}_i, \quad \tilde{x}_{i0} = (1, \circ, \circ), \quad i = 1, 2, \dots, n \quad (1.3)$$

که در آن $\tilde{y}_i$ ، i امین عدد فازی مشاهده شده برای متغیر پاسخ و $\tilde{x}_{ij}$ یک عضو از ماتریس متغیرهای ورودی با اعداد فازی LR است. همچنین $\tilde{\beta}_j$ ، j امین عضو از بردار $\tilde{\beta} = (\tilde{\beta}_0, \tilde{\beta}_1, \dots, \tilde{\beta}_p)^T$ به عنوان ضرایب رگرسیونی با اعداد فازی LR و $\tilde{\epsilon}_i$ بیانگر جمله خطا است. مدل برآورد شده را می‌توان به صورت ۲.۳ نمایش داد:

$$\hat{y}_i = \oplus_{j=0}^p (\hat{\beta}_j \otimes \tilde{x}_{ij}), \quad \tilde{x}_{i0} = (1, \circ, \circ), \quad i = 1, 2, \dots, n \quad (2.3)$$

ار آنجایی که در استنباط آماری، یکی از مهمترین معیارها برای تجزیه و تحلیل عملکرد یک برآوردگر MPE است، در اینجا ما با منقبض کردن برآورد اولیه ضرایب مدل رگرسیون به سمت

^۲Integrated fuzzy shrinkage

صفر یا یک نقطه هدف، قادر به کاهش MPE آن برآورد هستیم. جزئیات بیشتر در مورد انواع متداول استفاده از برآورد انقباضی در مدل رگرسیون در [۱۱۸]، [۱۱۹] به تفصیل اشاره شده است. در این بین، انقباض نوع استاین [۷۴]، یک تابع غیر خطی از فاصله توان دوم بین عناصر برآورد خام و مقدار صفر است. در اینجا، پس از اجرای روش انقباضی نوع استاین بر روی ضرایب رگرسیون فازی، به عنوان یک اصلاحی، روش پیشنهادی انقباضی فازی یکپارچه (IFS) را برای غلبه بر اشکالات روش پیشنهادی ارائه می‌دهیم.

با فرض $\hat{\beta} = (\hat{\beta}_0, \hat{\beta}_1, \dots, \hat{\beta}_p)$ به عنوان یک برآورد از پارامترهای مدل رگرسیونی فازی (تولید شده توسط هر روش)، $\hat{\beta}^S = (\hat{\beta}_0^S, \hat{\beta}_1^S, \dots, \hat{\beta}_p^S)$ را به عنوان برآورد انقباضی نوع استاین برای $\hat{\beta}$ تعریف می‌کنیم که در آن

$$\hat{\beta}_j^S = (\hat{m}_{\beta_j}^S, \hat{l}_{\beta_j}^S, \hat{r}_{\beta_j}^S)_{LR} = \mathbf{h}_j^\top \mathbf{G}_j, \quad j = 0, 1, \dots, p \quad (۳.۳)$$

که در آن

$$\mathbf{G}_j = \begin{pmatrix} \hat{m}_{\beta_j} & 0 \\ & \hat{l}_{\beta_j} \\ 0 & & \hat{r}_{\beta_j} \end{pmatrix} = \text{Diag}(\hat{m}_{\beta_j}, \hat{l}_{\beta_j}, \hat{r}_{\beta_j})^\top \quad (۴.۳)$$

و

$$\mathbf{H}_j = \left(\left(1 - \frac{k}{\hat{m}_{\beta_j}^\vee} \right), \left(1 - \frac{k}{\hat{l}_{\beta_j}^\vee} \right), \left(1 - \frac{k}{\hat{r}_{\beta_j}^\vee} \right) \right)^\top \quad (۵.۳)$$

در اینجا $k \in \mathbb{R}^+$ را ثابت انقباضی می‌نامیم که با تغییر آن میزان انقباض به سمت صفر تنظیم می‌شود. لازم به ذکر است که هزینه محاسباتی در این روش بسیار ناچیز است چرا که ما برای ساخت برآورد انقباضی از برآورد اولیه استفاده می‌کنیم و تا زمانی که کاهش MPE را بدنبال داشته باشد تغییر را اعمال می‌کنیم. برای فهم بهتر به مثال زیر توجه کنید.

به طور معمول برای نرم کردن بیسکوییت، آن را به صورت عمودی وارد لیوان چای می‌کنیم و هدف، رسیدن به بیشترین بخش نرم شده در بیسکوییت است. فرض کنید این یک ساختار برآوردی با خاصیت ناریبی باشد. اکنون با تغییر در حالت عمودی آن به اریب، با اینکه از حالت ناریبی خارج می‌شویم اما نرم شدن سطح بیستری از بیسکوییت را بدست می‌آوریم. این تغییر تا میزان خاصی قابل اعمال است و تغییر بیش‌تر از آن مجاز نیست. از این رو نه تنها در بازه‌ای از تغییر، هدف خود را بهبود داده‌ایم بلکه در زاویه خاصی از این تغییر، بیشترین بهینگی در هدف را خواهیم داشت. این همان فرآیندی است که توسط ثابت انقباضی در برآورد نوع استاین انجام می‌گیرد.

در بکارگیری این روش در مدل‌سازی رگرسیون خطی فازی^۳ (FLR) با وجود اینکه در مدل رگرسیون کلاسیک MPE کوچکتر را در برآورد فراهم می‌کند با دو مشکل مواجه خواهیم شد:

^۳Fuzzy linear regression

۱. از آنجایی که مقادیر پهنا همواره مثبت هستند، با افزایش مقدار k امکان منفی شدن آن‌ها وجود دارد.

۲. از طرف دیگر، برای هر کدام از مدل‌های مورد بررسی، بر اساس معیارهای ارزیابی مختلف، یک ثابت انقباضی بهینه خواهیم داشت. از این رو، در هر مدل FLR چندین مقدار مختلف بهینه k فراهم می‌شود. بدین معنی که علاوه بر تولید یک مقدار بهینه از k در هر مدل و بر اساس هر معیار ارزیابی ما در بازه‌ای از این مقادیر کماکان نسبت به مدل اولیه MPE بهتری داریم. در این حالت، محقق ممکن است نتواند در مورد یک k بهینه تصمیم بگیرد. برای غلبه بر این دو مشکل، روش انقباضی نوع استاین را به شرح زیر اصلاح می‌کنیم.

روش پیشنهادی IFS

برای رفع مشکل اول بجای $(\hat{l}_{\beta_j}^S)(\hat{r}_{\beta_j}^S)$ از $(\hat{l}_{\beta_j}^{S+})(\hat{r}_{\beta_j}^{S+})$ استفاده می‌کنیم که (آن را به نام برآورد استاین نوع مثبت می‌خوانیم) عبارتند از

$$\begin{cases} \hat{l}_{\beta_j}^{S+} = \left[\left(1 - \frac{k}{\hat{l}_{\beta_j}^S} \right) \hat{l}_{\beta_j} \right]^+ = \max \left\{ 0, \left(1 - \frac{k}{\hat{l}_{\beta_j}^S} \right) \hat{l}_{\beta_j} \right\} \\ \hat{r}_{\beta_j}^{S+} = \left[\left(1 - \frac{k}{\hat{r}_{\beta_j}^S} \right) \hat{r}_{\beta_j} \right]^+ = \max \left\{ 0, \left(1 - \frac{k}{\hat{r}_{\beta_j}^S} \right) \hat{r}_{\beta_j} \right\} \end{cases} \quad (6.3)$$

همچنین برای رفع مشکل دوم، ز یک روش میانگین‌سازی مدل استفاده می‌کنیم که در آن از بین تمام بازه‌های بهینه انقباض، حداکثر مقدار اشتراک بازه ثابت‌های انقباضی را به عنوان ثابت انقباضی میانگین استفاده می‌کنیم. به بیان دقیق تر، از آنجا که در بهینه‌سازی هر معیار ارزیابی GOF، یک محدوده مطلوب برای مقادیر k پیدا می‌کنیم، پس از بدست آوردن اشتراک این بازه‌ها، حداکثر مقدار بازه اشتراکی را به عنوان ثابت انقباضی مشترک در نظرمی‌گیریم. بنابراین، رویکرد انقباضی یکپارچه IFS به عنوان مدل میانه رو، نقطه قوت کار و نوآوری است. نتایج حاصل از این عملیات را می‌توانید در مثال‌های شبیه‌سازی و واقعی که در بخش ۳.۳ ارائه خواهد شد ببینید.

۲.۳ آزمون فرضیه فازی

برای ارزیابی عملکرد روش پیشنهادی در مدل‌های چند متغیره، در این بخش، روش آزمایش فرضیه فازی که توسط لی و همکاران در سال ۲۰۱۵ [۸۵] و بر اساس یک نمونه بوت استرپ پیشنهاد شده است را ارائه می‌دهیم. فرضیه فازی را به صورت زیر در نظر بگیرید:

$$H_0 : \tilde{\beta}_{j_\alpha} = 0, \quad j = 0, 1, \dots, p \quad (7.3)$$

که در آن $\tilde{\beta}_{j\alpha}$ ، آلفا-برش j امین ضریب مدل رگرسیونی فازی است به صورت $\tilde{\beta}_{j\alpha} = [\tilde{\beta}_{j\alpha}^L, \tilde{\beta}_{j\alpha}^U]$ آنگاه با در نظر گرفتن $\beta_{j\alpha}^C = \frac{1}{2}(\tilde{\beta}_{j\alpha}^U + \tilde{\beta}_{j\alpha}^L)$ و $\beta_{j\alpha}^W = \frac{1}{2}(\tilde{\beta}_{j\alpha}^U - \tilde{\beta}_{j\alpha}^L)$ که در آن $\beta_{j\alpha}^W \geq 0$ است تعریف می کنیم

$$d^2(\tilde{\beta}_{j\alpha}, 0) = \sqrt{(\beta_{j\alpha}^C)^2 + (\beta_{j\alpha}^W)^2}$$

از اینرو آماره آزمون بکار رفته عبارتست از:

$$T_\alpha = \sqrt{d^2\left(\frac{\hat{\beta}_{j\alpha}}{s_{\hat{\beta}_{j\alpha}}}, 0\right)} \quad (۸.۳)$$

که در آن $s_{\hat{\beta}_{j\alpha}} = \sqrt{\frac{\sum_{i=1}^n d^2(\hat{y}_{i\alpha}, \bar{y}_{i\alpha})}{(n-p-1) \sum_{i=1}^n (x_{ij} - \bar{x}_j)^2}}$ انحراف معیار $\hat{\beta}_{j\alpha}$ است. اکنون برای اجرای آزمون و بررسی رد یا پذیرش فرض صفر در هر آلفا-برش بر اساس یک نمونه بوت استرپ، الگوریتم ۳ را در نظر می گیریم.

الگوریتم ۳ - الگوریتم محاسباتی برای P - مقدار در آزمون (۷.۳)

گام اول تا پنجم را تا زمانی که $\alpha_h = 1$ که در آن $\alpha_h = 1, \dots, \Delta\alpha \rightarrow 0$ ، $h = 0, \dots, \Delta\alpha$ است ادامه دهید. مقدار اولیه برای α_0 صفر است.

گام ۱: تعداد B^* نمونه بوت استرپ در نظر بگیرید.

گام ۲: ضرایب مدل رگرسیونی فازی را برای نمونه های بوت استرپ محاسبه کرده (جزئیات را در [۴۸، ۸۵] ببینید). و مقادیر انقباضی متناظر با هر کدام را با استفاده از بهترین مقدار پارامتر تنظیم در مدل IFS از طریق معادلات (۳.۳) تا (۶.۳) بدست آورید و آن را با $\hat{\beta}_j^{(b^*)}$ ، $b^* = 1, 2, \dots, B^*$ نمادگذاری کنید.

گام ۳: توزیع تجربی آماره آزمون را بر اساس مقادیر $T_\alpha^{(b^*)}$ مشخص کنید [۸۵].

گام ۴: عدد p - مقدار را از طریق نسبت مقادیری از گروه $\{T_\alpha^{(1)}, T_\alpha^{(2)}, \dots, T_\alpha^{(B^*)}\}$ که بزرگتر یا مساوی T_α هستند در نظر بگیرید.

گام ۵: فرض صفر را در سطح معنی داری γ رد کنید هرگاه

$$p - value = \frac{\sum_{b^*=1}^{B^*} I(T_\alpha^{(b^*)} \geq T_\alpha)}{B^*} < \gamma$$

که در آن $I(\cdot)$ بیانگر تابع نشانگر است. یعنی هرگاه شرط داخل پرانتز در فرمول بالا برقرار باشد یک است و در غیر اینصورت صفر.

۳.۳ مطالعات عددی

در این بخش، برای مقایسه عملکرد روش پیشنهادی IFS با برخی از روش‌های موجود در این زمینه، مثال‌هایی را در دو بخش مثال شبیه‌سازی و واقعی ارائه کرده و علاوه بر نمایش عملکرد بهینه مدل پیشنهادی در مقابل دیگر مدل‌های مورد آزمایش طبق معیارهای بهینگی مدل GOF، آزمون فرضیه متناظر با ضرایب مدل را نیز اجرا و به نمایش می‌گذاریم.

۱.۳.۳ مطالعه شبیه‌سازی

در این مثال شبیه‌سازی ما روش پیشنهادی IFS را بر روی مدل به‌دست آمده توسط کیم و بیشو [۸۲] که در سال ۱۹۹۸ ارائه شده‌است، اجرا کرده‌ایم و آن را تحت معیارهای GOF با مدل‌های [۸۱] (KT) و [۱۵۰] (Zeng) مقایسه می‌کنیم. از این رو با مشخصات زیر یک نمونه ۳۰ تایی را از مدل (۱.۳) تولید و فرایند را ۳۰ بار تکرار می‌کنیم.

$$x_{i1} = z_1 + \epsilon_1, z_1 \sim N(3, 1/5) \quad , \quad \epsilon_1 \sim N(0, 0/05) \quad ۱.$$

$$x_{i2} = z_2 + \epsilon_2, z_2 \sim N(2, 1) \quad , \quad \epsilon_2 \sim N(0, 0/06)$$

$$x_{i3} = z_3 + \epsilon_3, z_3 \sim N(0, 1) \quad , \quad \epsilon_3 \sim N(0, 0/04)$$

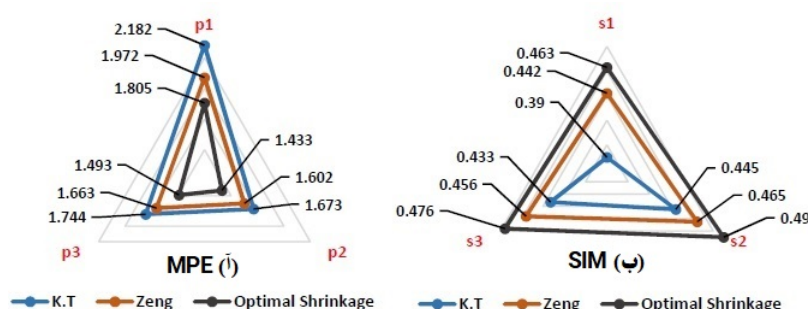
$$\tilde{\beta}_0 = (3, 0/2)_T, \quad \tilde{\beta}_1 = (0/11, 1)_T, \quad \tilde{\beta}_2 = (1/5, 0/5)_T, \quad \tilde{\beta}_3 = (0, 0/2)_T \quad ۲.$$

$$\tilde{\epsilon}_i = (\epsilon_i, s_{\epsilon_i})_T, s_{\epsilon_i} = |\epsilon_i| \quad , \quad \epsilon_i \sim N(0, 2) \quad ۳.$$

نتایج ارزیابی برازش سه مدل انقباضی، KT و Zeng در جدول ۱.۳ به طور خلاصه ارائه شده است. از آنجایی که MPE بیانگر متوسط خطای پیشگویی و SIM به مفهوم میزان شباهت پاسخ برآورد شده با مشاهده شده است لذا کاهش در MPE و افزایش در SIM برتری مدل را نشان می‌دهد. بنابراین، همانطور که مشاهده می‌شود مدل پیشنهادی IFS دارای مقادیری کوچکتر در معیارهای MPE و مقادیری بزرگتر در معیارهای SIM نسبت به دیگر مدل‌های مورد بررسی است و این نشان از برتری مدل پیشنهادی است.

جدول ۱.۳: مقادیر GOF برای ارزیابی میانگین ۳۰ مجموعه داده شبیه‌سازی شده برای روش انقباضی در مقابل مدل‌های دیگر.

معیارهای GOF						مدل
S_3	S_2	S_1	P_3	P_2	P_1	
۰/۴۷۶	۰/۴۹۰	۰/۴۶۳	۱/۴۹۳	۱/۴۳۳	۱/۸۰۵	انقباضی
۰/۴۳۳	۰/۴۴۵	۰/۳۹۰	۱/۷۴۴	۱/۶۷۳	۲/۱۸۲	KT
۰/۴۵۶	۰/۴۶۵	۰/۴۴۲	۱/۶۶۳	۱/۶۰۲	۱/۹۷۲	Zeng



شکل ۱.۳: مقایسه شهودی سه مدل بر اساس نتایج داده‌های شبیه‌سازی: (آ) مقادیر MPE (ب) مقادیر SIM.

همچنین با توجه به نمودار ۱.۳ سه گانه مربوط به معیار MPE در خصوص مدل پیشنهادی (مثلت داخلی) در درون دو مثلث دیگر که برای دو مدل رقیب طراحی شده است قرار دارد که بیانگر اثر بخشی مدل پیشنهادی در کاهش معیار MPE است. این اثر بخشی در نمودار (ب) با ارائه سه گانه افزایش یافته در معیار شباهت SIM برای مدل پیشنهادی در مقابل مدل‌های رقیب قابل مشاهده است. این نکته قابل تامل است که ما روش انقباضی را بر روی مدلی مربوط به سال ۱۹۹۸ [۸۲] پیاده کرده‌ایم و آن را با مدل‌های مربوط به سال‌های ۲۰۱۲ [۸۱] و ۲۰۱۷ [۱۵۰] مقایسه کرده‌ایم. این برتری نشان می‌دهد که با استفاده از انقباض در مدل‌های قدیمی هنوز هم ممکن است عملکرد خوبی در مقایسه با مدل‌های جدید به دست آوریم. ضرایب مربوط به بهترین مدل در این فرآیند را که از ۱۸ امین تکرار مدل شبیه‌سازی شده به دست آمده است در جدول ۲.۳ به همراه مقادیر GOF متناظر آن ارائه کرده‌ایم.

جدول ۲.۳: نتایج ارزیابی روش انقباضی در ۱۸ امین مجموعه داده شبیه‌سازی شده.

ضرایب	P_1	P_2	P_3	S_1	S_2	S_3
$\hat{\beta}_0 = (3/515, 0/203)_T$	۱/۲۲۴	۱/۲۷۹	۱/۵۵۲	۰/۵۰۶	۰/۴۹۰	۰/۴۷۵
$\hat{\beta}_1 = (0/118, 0/999)_T$						
$\hat{\beta}_2 = (0/790, 0/504)_T$						
$\hat{\beta}_3 = (0/028, 0/103)_T$						

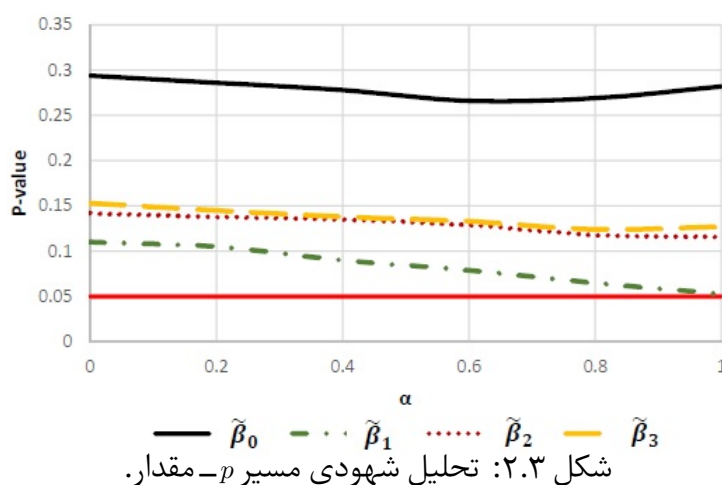
نزدیکی تقریبی مقادیر برآورد شده با مقادیر واقعی پارامترها نشان از عملکرد این روش در دستیابی به یک مدل مناسب است.

اکنون و در ادامه فرآیند نمایش بهینگی مدل به بررسی آزمون فرضیه ضرایب بر اساس فرض ۹.۳ می‌پردازیم:

$$H_0 : \tilde{\beta}_{j\alpha} - \tilde{\beta}_{j\alpha}^* = 0, \quad j = 0, 1, 2, 3 \quad (9.3)$$

که در آن $\tilde{\beta}_{j\alpha}^*$ ، ها مقادیر واقعی پارامترها در مدل شبیه‌سازی شده اند. نتایج آزمون مربوطه در سطح معنی داری ۵٪ در جدول ۳.۳ آمده است. علاوه بر آن شکل ۲.۳، رفتار P - مقدارهای

به‌دست آمده در آلفا-برش‌های مختلف را نشان می‌دهد.



جدول ۳.۳: نتایج آزمون فرضیه در داده‌های شبیه‌سازی شده ۱۸ ام برای $\tilde{\beta}_{j\alpha}$ همراه با T_α و معیار p -مقدار.

آلفا - برش						ضریب
۱/۰	۰/۸	۰/۶	۰/۴	۰/۲	۰/۰	
۱/۸۸۷	۱/۹۱۲	۱/۹۷۷	۲/۱۰۸	۲/۲۱۴	۲/۲۲۵	T_α
۰/۲۸۲	۰/۲۶۹	۰/۲۶۶	۰/۲۷۸	۰/۲۸۶	۰/۲۹۴	p -مقدار $\tilde{\beta}_{0\alpha}$
۱/۵۱۷	۱/۷۶۶	۱/۸۴۳	۱/۹۸۵	۲/۲۱۶	۲/۳۰۱	T_α
۰/۰۵۳	۰/۰۶۵	۰/۰۷۹	۰/۰۹۱	۰/۱۰۵	۰/۱۱۱	p -مقدار $\tilde{\beta}_{1\alpha}$
۱/۵۱۷	۱/۷۶۶	۱/۸۴۳	۱/۹۸۵	۲/۲۱۶	۲/۳۰۱	T_α
۰/۱۱۶	۰/۱۱۱۸	۰/۱۲۹	۰/۱۳۵	۰/۱۳۸	۰/۱۴۲	p -مقدار $\tilde{\beta}_{2\alpha}$
۱/۵۱۷	۱/۷۶۶	۱/۸۴۳	۱/۹۸۵	۲/۲۱۶	۲/۳۰۱	T_α
۰/۱۲۷	۰/۱۲۴	۰/۱۳۳	۰/۱۳۸	۰/۱۴۵	۰/۱۵۳	p -مقدار $\tilde{\beta}_{3\alpha}$

بر اساس نتایج حاصله که در جدول ۳.۳ و شکل ۲.۳ مشاهده می‌شود از آنجایی که p -مقدارهای حاصله در جدول و نمایش مسیر حرکت آنها در آلفا-برش‌های مختلف در شکل فوق، همگی بالاتر از سطح خطای ۵٪ می‌باشد لذا فرضیه صفر در کل آزمون‌های مربوط به ضرایب در سطح معنی داری ۵٪ پذیرفته می‌شود. بدین معنی که تمامی ضرایب برآورد شده در مدل مانده و هیچکدام را نمی‌توان حذف نمود.

۲.۳.۳ تحلیل داده‌های واقعی

در این بخش، عملکرد روش انقباضی را در داده‌های واقعی و نسبت به مدل‌های اولیه و بدون تغییر در شرایط و تابع هدف مدل مبنا، با استفاده از معیارهای ارزیابی GOF مورد بررسی قرار می‌دهیم. همچنین وجود یا عدم وجود ضرایب مدل IFS طی یک آزمون فرضیه فازی تحلیل می‌گردد. برای تکمیل اثر بخشی مدل‌های پیشنهادی که در بهینه‌ترین حالت ممکن برای معیارهای MPE در مقابل مدل پایه در مثال‌های عددی تولید شده‌اند از نمودار شهودی-استنباطی تیلور که در بخش ۴.۱ تشریح شده استفاده کرده‌ایم.

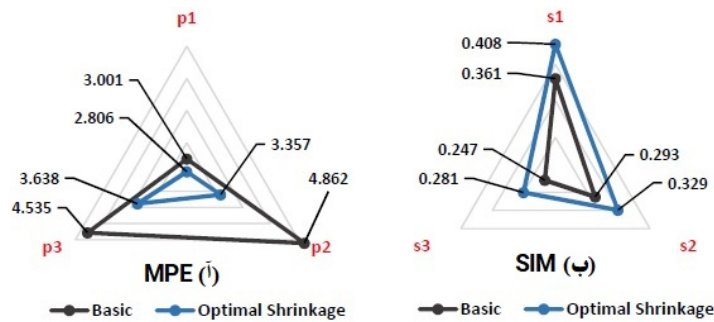
مثال اول

در این مثال از یک مجموعه داده واقعی که در بخش مثال‌های انگیزشی در ۱.۲ به عنوان اولین مجموعه داده آمده، استفاده می‌کنیم. بر اساس این داده‌ها کلکین‌نما و طاهری در [۸۱] نشان دادند که مدل آن‌ها در مقایسه با مدل چوی و باکلی [۳۴] عملکرد بهتری دارد. در این مثال می‌خواهیم مدل انقباضی را با مدل ارائه شده در [۸۱] مقایسه کنیم. نتایج ارزیابی، پس از تولید مدل‌های مورد بحث، توسط معیارهای ارزیابی GOF در جدول ۴.۳ به طور خلاصه ارائه شده‌اند.

جدول ۴.۳: مقایسه مدل انقباضی و مدل پایه بر اساس معیارهای ارزیابی GOF در مثال اول.

معیارهای ارزیابی GOF						مدل
S_3	S_2	S_1	P_3	P_2	P_1	
۰/۲۴۷	۰/۲۹۳	۰/۳۶۱	۴/۵۳۳	۴/۸۶۲	۳/۰۰۱	مبنا [۸۱]
۰/۲۸۱	۰/۳۲۹	۰/۴۰۸	۳/۶۳۸	۳/۳۵۷	۲/۸۰۶	انقباضی
۰/۰۴۶	۰/۰۴۹	۰/۰۵۰	۰/۰۵۴	۰/۰۴۰	۰/۰۴۹	ثابت انقباضی بهینه k
(۰, ۰/۱۸۳]	(۰, ۰/۱۹۹]	(۰, ۰/۱۷۳]	(۰, ۰/۱۶۷]	(۰, ۰/۱۸۸]	(۰, ۰/۱۳۶]	محدوده بهینگی

در این جدول همانطور که قبلاً اشاره شده بود مدل مبنا با استفاده از روش پیشنهاد شده در [۸۱] تولید شده است و مدل انقباضی ارائه شده نیز بر اساس روش انقباضی نوع مثبت استاین (برای رفع مشکل نامنفی بودن پهناها) به‌دست آمده است. بر اساس نتایج به‌دست آمده در جدول ۴.۳ برای هر معیار ارزیابی GOF، مقادیر بهینه k متناظر با آن‌ها را (اعداد داخل پرانتز) که در آن بیشترین برتری در بهینگی مدل را نسبت به مدل مبنا داشته ارائه شده است. علاوه بر آن براساس هر معیار بهینگی GOF، بازه‌هایی که در آن مدل انقباضی کماکان از مدل مبنا عملکرد بهتری دارد (متناظر با هر معیار GOF) با عنوان محدوده بهینگی، ارائه شده است. این برتری در شکل ۳.۳ به نمایش گذاشته شده است.



شکل ۳.۳: مقایسه شهودی مدل انقباضی و مدل پایه بر اساس معیارهای ارزیابی در مثال اول: (ا) مقدار MPE (ب) مقدار SIM.

اکنون با در نظر گرفتن آزمون فازی زیر به بررسی و آزمون ضرایب مدل برای مدل بدست آمده در IFS می‌پردازیم:

$$H_0: \tilde{\beta}_{j\alpha} = 0, \quad j = 0, 1, 2, 3 \quad (10.3)$$

نتایج آزمون فرض فوق برای هر $j \in \{0, 1, 2, 3\}$ در جدول ۵.۳ خلاصه شده است.

جدول ۵.۳: آزمون فرض هر $\tilde{\beta}_{j\alpha}$ همراه با T_α و معیار p - مقدار.

آلفا - برش						ضریب
۱/۰	۰/۸	۰/۶	۰/۴	۰/۲	۰/۰	
۲/۷۱۷	۲/۸۳۰	۲/۹۶۱	۳/۰۸۳	۳/۳۱۶	۳/۴۶۲	T_α
۰/۰۱۸	۰/۰۲۳	۰/۰۲۱	۰/۰۴۲	۰/۰۴۴	۰/۰۴۷	$\tilde{\beta}_{0\alpha}$ مقدار - p
۲/۸۶۷	۲/۹۲۲	۲/۹۵۳	۳/۰۸۵	۳/۲۵۳	۳/۳۶۷	T_α
۰/۰۱۹	۰/۰۲۷	۰/۰۳۲	۰/۰۴۳	۰/۰۴۰	۰/۰۴۶	$\tilde{\beta}_{1\alpha}$ مقدار - p
۲/۴۸۴	۲/۵۹۶	۲/۷۴۳	۲/۸۸۵	۲/۹۳۷	۳/۱۰۲	T_α
۰/۰۲۲	۰/۰۲۹	۰/۰۳۹	۰/۰۳۶	۰/۰۴۵	۰/۰۴۷	$\tilde{\beta}_{2\alpha}$ مقدار - p
۲/۵۱۷	۲/۷۶۶	۲/۸۴۳	۲/۹۸۵	۳/۰۱۱	۳/۱۳۰	T_α
۰/۰۱۹	۰/۰۳۳	۰/۰۳۶	۰/۰۴۲	۰/۰۴۴	۰/۰۴۸	$\tilde{\beta}_{3\alpha}$ مقدار - p

در جدول ۶.۳ مدل‌های بدست آمده برای روش مبنا و همچنین مدل‌های انقباضی متناظر با مقدار بهینه هر شش معیار بهینگی به همراه مدل میانگین IFS ارائه شده است. توجه داشته باشید که بر اساس تعریف مدل IFS، مقدار ثابت انقباضی در این مدل بر اساس ماکزیمم مقدار بازه بدست آمده از اشتراک تمام بازه‌های بهینگی حاصل می‌گردد که در این مثال با توجه به بازه‌های تولیدی، مقدار عددی ثابت انقباضی برابر با ۰/۱۳۶ خواهد بود.

جدول ۶.۳: مدل‌های برازش شده مبنا، بهینه انقباضی و IFS برای مثال اول.

مدل	روش
$\tilde{y}_i = (-14/900, 0/250)_T \oplus_w (-0/251, 0/250)_T \odot_w x_{i1}$ $\oplus_w (-0/956, 0/222)_T \odot_w x_{i2} \oplus_w (1/467, 0/084)_T \odot_w x_{i3}$	مبنا [۸۱]
$\tilde{y}_i = (-14/896, 0/034)_T \oplus_w (-0/035, 0/034)_T \odot_w x_{i1}$ $\oplus_w (-0/899, 0/000)_T \odot_w x_{i2} \oplus_w (1/427, 0/000)_T \odot_w x_{i3}$	۱
$\tilde{y}_i = (-14/897, 0/090)_T \oplus_w (-0/091, 0/090)_T \odot_w x_{i1}$ $\oplus_w (-0/914, 0/041)_T \odot_w x_{i2} \oplus_w (1/437, 0/000)_T \odot_w x_{i3}$	۲
$\tilde{y}_i = (-14/896, 0/054)_T \oplus_w (-0/055, 0/054)_T \odot_w x_{i1}$ $\oplus_w (-0/904, 0/000)_T \odot_w x_{i2} \oplus_w (1/431, 0/000)_T \odot_w x_{i3}$	۳
$\tilde{y}_i = (-14/896, 0/054)_T \oplus_w (-0/055, 0/054)_T \odot_w x_{i1}$ $\oplus_w (-0/904, 0/000)_T \odot_w x_{i2} \oplus_w (1/431, 0/000)_T \odot_w x_{i3}$	۴
$\tilde{y}_i = (-14/897, 0/066)_T \oplus_w (-0/067, 0/066)_T \odot_w x_{i1}$ $\oplus_w (-0/908, 0/014)_T \odot_w x_{i2} \oplus_w (1/433, 0/000)_T \odot_w x_{i3}$	۵
$\tilde{y}_i = (-14/896, 0/050)_T \oplus_w (-0/051, 0/050)_T \odot_w x_{i1}$ $\oplus_w (-0/903, 0/000)_T \odot_w x_{i2} \oplus_w (1/430, 0/000)_T \odot_w x_{i3}$	۶
$\tilde{y}_i = (-14/891, 0/000)_T \oplus_w (0/292, 0/000)_T \odot_w x_{i1}$ $\oplus_w (-0/814, 0/000)_T \odot_w x_{i2} \oplus_w (1/371, 0/000)_T \odot_w x_{i3}$	IFS

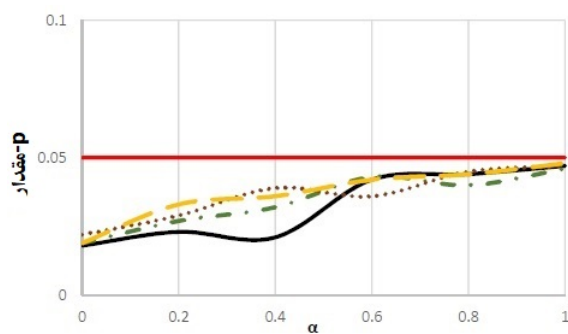
تذکر ۱.۳.۳. از آنجایی که در بکارگیری رویکرد IFS، بهینگی مدل را در تمامی معیارهای ارزیابی مدنظر داریم لذا این باعث فشردگی در بازه های بهینگی شده و ممکن است اثر آن در صفر شدن پهنای مدل برازش شده تحت رویکرد IFS ظاهر گردد. که این مسئله بیشتر زمانی رخ می‌دهد که پهنای مدل مبنا به لحاظ مقدار عددی کوچک باشند. لذا در این حالت پیشنهاد می‌شود که از میانگین مقادیر بهینه k برای مدل IFS استفاده شود.

بنابراین مدل IFS در این مثال با $k = 0/47$ به صورت است.

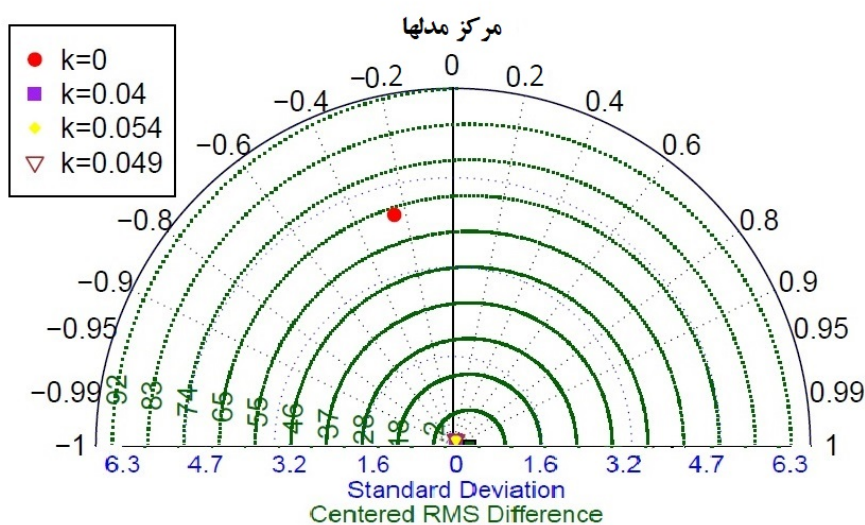
$$\tilde{y}_i = (-14/897, 0/057)_T \oplus_w (-0/058, 0/060)_T \odot_w x_{i1}$$

$$\oplus_w (-0/905, 0/01)_T \odot_w x_{i2} \oplus_w (1/405, 0/000)_T \odot_w x_{i3}$$

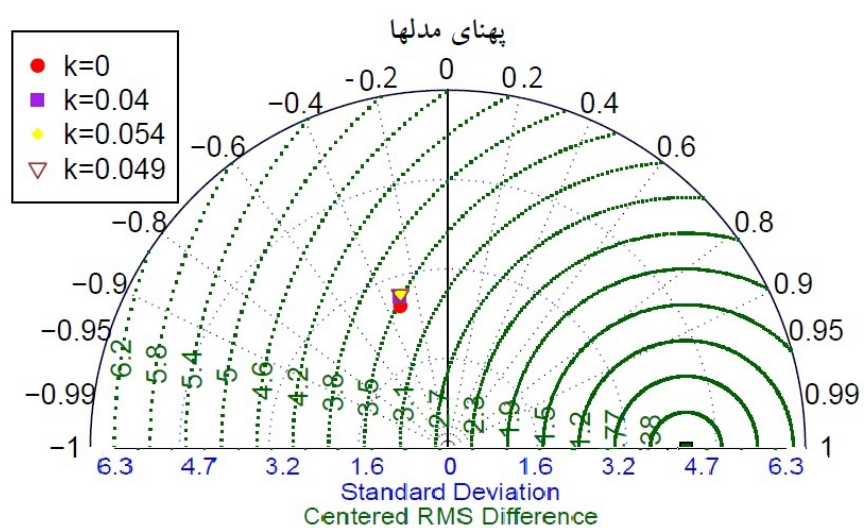
اکنون بر اساس p -مقدارهای بدست آمده برای تمامی ضرایب در آلفا-برش‌های مختلف می‌توان در سطح معنی داری ۵٪ فرضیه صفر را در خصوص پارامترهای مدل IFS رد کرد. شکل ۴.۳ این نتیجه را به صورت شهودی نشان می‌دهد.



شکل ۴.۳: تحلیل شهودی مسیر معیار p -مقدار.



(ا)



(ب)

شکل ۵.۳: مقایسه شهودی و استنباطی مدل پایه و مدل‌های انقباضی در مقادیر مختلف k ، در مثال اول. (ا) مقدار مرکز (ب) مقدار پهنای.

در اینجا بر اساس ساختار نمودار تیلور که در ۴.آ ارائه شده است، اثر بخشی مدل‌های انقباضی نوع استاتین پیشنهادی را در مقایسه با مدل مبنا در نمودار ۵.۳ برای هر دو بخش عدد فازی (مرکز و پهنا) نشان داده‌ایم. این نمودار همزمان سه مولفه را مورد مقایسه قرار می‌دهد. اولین مولفه میزان همبستگی پاسخ‌های برآورد شده با پاسخ‌های مشاهده شده است که مقادیر همبستگی در روی محیط دایره قابل انطباق است. مولفه دوم انحراف معیار پاسخ‌های برآورد شده است که نیم دایره‌هایی به شعاع مقادیر داده شده در محور افقی و مشخص شده با نقطه چین، این مولفه را پوشش می‌دهند. مولفه سوم هم خطای استاندارد تولید شده توسط پاسخ‌های برآورد شده نسبت به مقادیر مشاهده شده است که این اعداد مربوط به این اختلاف را می‌توان با مقادیر داده شده روی خط مورب در شکل ۵.۳ رصد کرد.

طبق نمودار ۵.۳، در بخش (آ)، میزان همبستگی مدل‌های پیشنهادی که با مثلث و نزدیک به مرکز خط افقی قرار دارد، بواسطه قرار گیری در خط مورب منتهی شده به مقدار عددی ۰/۹، با مشاهدات واقعی حدود ۹۰٪ است که این میزان برای مدل مبنا که با دایره توپر نشان داده شده است حدودا ۰/۲۵- نشان داده شده است. همچنین انحراف معیار و خطای استاندارد تولید شده توسط مدل‌های پیشنهادی به دلیل قرار گیری در نیم دایره کوچکتر، بسیار کمتر از مدل مبنا است. در بخش (ب) نیز (پهنای مدل) اختلاف قابل ملاحظه‌ای نسبت به مدل مبنا دیده نمی‌شود.

مثال دوم

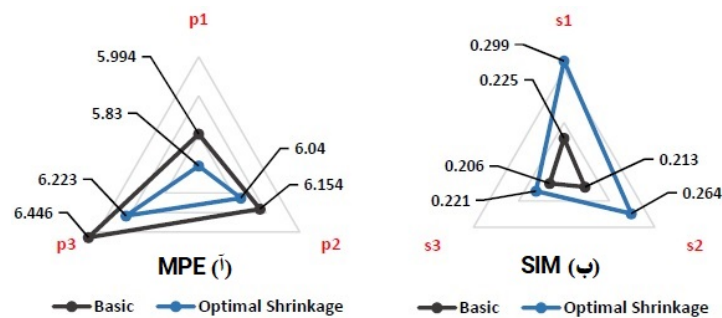
داده‌های این مثال در بخش مثال‌های انگیزشی به عنوان دومین مجموعه داده در ۲.۲ ارائه شده است. در این مثال از مدل ارائه شده توسط حسن پور و همکاران [۶۰] در سال ۲۰۰۹ به عنوان مدل مبنا استفاده کرده‌ایم.

جدول ۷.۳: مقادیر GOF برای ارزیابی مدل‌های پایه انقباضی در مثال دوم.

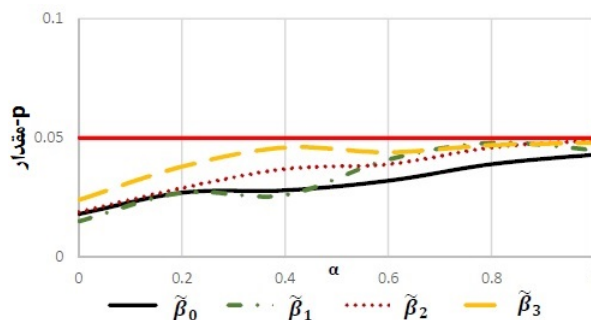
معیارهای GOF						مدل
S_3	S_2	S_1	P_3	P_2	P_1	
۰/۲۰۶	۰/۲۱۳	۰/۲۲۵	۶/۴۴۶	۶/۱۵۴	۵/۹۹۴	مبنا [۶۰]
۰/۲۲۱	۰/۲۴۶	۰/۲۹۹	۶/۲۲۳	۶/۰۴	۵/۸۳۰	انقباضی
۲/۳۹۴	۱/۵۱۰	۱/۴۱۸	۱/۸۶۰	۱/۱۰۰	۱/۱۸۳	ثابت انقباضی بهینه (k)
[۰،۳/۲۴۵)	[۰،۲/۳۸۰)	[۰،۲/۲۴۶)	[۰،۳/۰۶۰)	[۰،۱/۷۳۰)	[۰،۱/۷۵۹)	محدوده بهینگی

از آنجایی که با کاهش در مقادیر MPE و افزایش در مقادیر SIM نسبت به دیگر مدل‌ها، بهینگی مدل را نتیجه خواهیم گرفت لذا با توجه به نتایج حاصل از جدول ۷.۳ در بخش معیارهای بهینگی GOF، کاهش مقادیر MPE و افزایش مقادیر SIM برای مدل پیشنهادی، برتری مدل به دست آمده با روش انقباضی نوع استاتین در مقابل مبنا را نتیجه می‌گیریم. مقادیر

k بهینه متناظر با هر کدام از معیارها و محدوده بهینگی آن‌ها نیز در این جدول ارائه شده است. گزارش شهودی این برتری در نمودار ۶.۳ قابل مشاهده است.



شکل ۶.۳: مقایسه شهودی مدل پایه و انقباضی تحت معیارهای ارزیابی. (آ) مقدار MPE (ب) مقدار SIM.



شکل ۷.۳: تحلیل شهودی مسیر معیار p - مقدار.

جدول ۸.۳: آزمون فرض هر ضریب $\tilde{\beta}_{j\alpha}$ همراه با T_α و معیار p - مقدار.

آلفا - برش						ضریب
۱/۰	۰/۸	۰/۶	۰/۴	۰/۲	۰/۰	
۲/۸۰۷	۲/۹۸۰	۳/۱۹۶	۳/۳۸۳	۳/۵۰۸	۳/۶۱۲	T_α
۰/۰۱۸	۰/۰۲۷	۰/۰۲۸	۰/۰۳۲	۰/۰۳۹	۰/۰۴۳	$\tilde{\beta}_{0\alpha}$ مقدار- p
۲/۹۶۷	۲/۹۷۱	۳/۰۰۵	۳/۱۷۲	۳/۲۷۴	۳/۴۱۵	T_α
۰/۰۱۵	۰/۰۲۷	۰/۰۲۶	۰/۰۴۱	۰/۰۴۸	۰/۰۴۵	$\tilde{\beta}_{1\alpha}$ مقدار- p
۲/۸۸۴	۲/۹۳۶	۲/۹۹۳	۳/۰۸۵	۳/۱۳۷	۳/۳۸۶	T_α
۰/۰۱۹	۰/۰۲۹	۰/۰۳۷	۰/۰۳۹	۰/۰۴۶	۰/۰۴۹	$\tilde{\beta}_{2\alpha}$ مقدار- p
۲/۶۴۱	۲/۷۳۰	۲/۸۶۵	۲/۹۰۴	۲/۹۲۱	۳/۰۰۸	T_α
۰/۰۲۴	۰/۰۳۸	۰/۰۴۶	۰/۰۴۴	۰/۰۴۷	۰/۰۴۸	$\tilde{\beta}_{3\alpha}$ مقدار- p

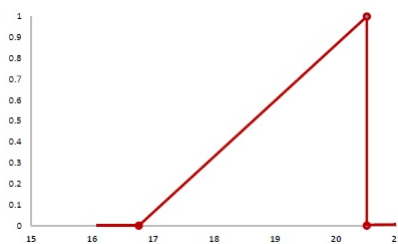
اکنون برای بررسی آزمون فرضیه ضرایب مدل پیشنهادی IFS در سطح معنی داری ۵٪، کافی است به نتایج به دست آمده در جدول ۸.۳ و شکل ۷.۳ مراجعه کنیم.

بر اساس p -مقدارهای بدست آمده برای تمامی ضرایب در آلفا-برش‌های مختلف چون این مقادیر کمتر از سطح معنی داری ۵٪ هستند، فرضیه صفر را در خصوص پارامترهای مدل IFS می‌توان رد کرد. بنابراین در سطح معنی داری ۵٪ وجود ضرایب به‌دست آمده در مدل پیشنهادی تایید می‌گردد.

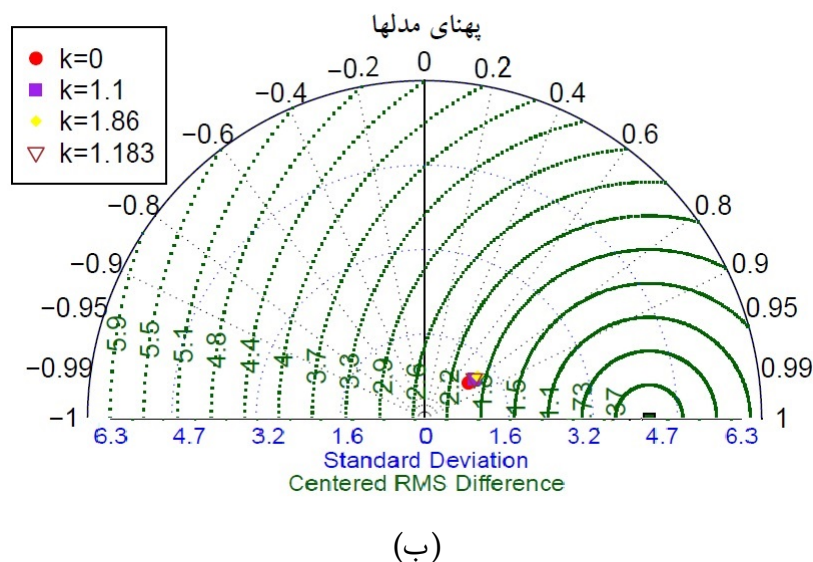
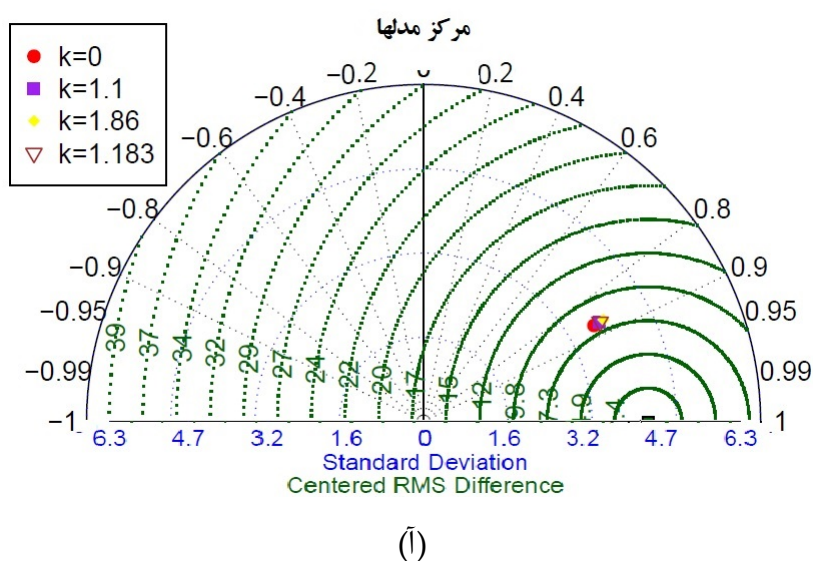
جدول ۹.۳: مدل‌های برازش شده برای هر ثابت انقباضی بهینه در مثال دوم.

مدل	روش
$\tilde{y}_i = (-127/693, 0/00, 0/00)_T \oplus (31/115, 0/00, 0/00)_T \otimes x_{i1}$ $\oplus (-2/919, 0/573, 0/00)_T \otimes x_{i2} \oplus (2/764, 0/801, 0/00)_T \otimes x_{i3}$	پایه [۶۰]
$\tilde{y}_i = (-127/684, 0/00, 0/00)_T \oplus (31/079, 0/00, 0/00)_T \otimes x_{i1}$ $\oplus (-2/542, 0/635, 0/00)_T \otimes x_{i2} \oplus (2/366, 0/963, 0/00)_T \otimes x_{i3}$	۱
$\tilde{y}_i = (-127/678, 0/00, 0/00)_T \oplus (31/055, 0/00, 0/00)_T \otimes x_{i1}$ $\oplus (-2/282, 0/678, 0/00)_T \otimes x_{i2} \oplus (2/092, 1/076, 0/00)_T \otimes x_{i3}$	۲
$\tilde{y}_i = (-127/684, 0/00, 0/00)_T \oplus (31/077, 0/00, 0/00)_T \otimes x_{i1}$ $\oplus (-2/514, 0/640, 0/00)_T \otimes x_{i2} \oplus (2/336, 1/976, 0/00)_T \otimes x_{i3}$	۳
$\tilde{y}_i = (-127/681, 0/00, 0/00)_T \oplus (31/066, 0/00, 0/00)_T \otimes x_{i1}$ $\oplus (-2/402, 0/658, 0/00)_T \otimes x_{i2} \oplus (2/218, 1/024, 0/00)_T \otimes x_{i3}$	۴
$\tilde{y}_i = (-127/674, 0/00, 0/00)_T \oplus (31/038, 0/00, 0/00)_T \otimes x_{i1}$ $\oplus (-2/099, 0/708, 0/00)_T \otimes x_{i2} \oplus (1/898, 0/743, 0/00)_T \otimes x_{i3}$	۵
$\tilde{y}_i = (-127/682, 0/00, 0/00)_T \oplus (31/070, 0/00, 0/00)_T \otimes x_{i1}$ $\oplus (-2/433, 0/653, 0/00)_T \otimes x_{i2} \oplus (2/251, 1/010, 0/00)_T \otimes x_{i3}$	۶
$\tilde{y}_i = (-127/679, 0/00, 0/00)_T \oplus (31/060, 0/00, 0/00)_T \otimes x_{i1}$ $\oplus (-2/326, 0/671, 0/00)_T \otimes x_{i2} \oplus (2/138, 1/056, 0/00)_T \otimes x_{i3}$	IFS

مدل‌های برازش شده توسط مدل مبنا، مدل‌های انقباضی بهینه به‌دست آمده بر اساس معیارهای شش گانه ارزیابی و مدل پیشنهادی IFS در جدول ۹.۳ ارائه شده است. توجه داشته باشید که بر اساس تعریف مدل IFS، مقدار ثابت انقباضی در این مدل بر اساس ماکزیمم مقدار بازه به‌دست آمده از اشتراک تمام بازه‌های بهینگی حاصل می‌گردد که در این مثال با توجه به بازه‌های تولیدی، مقدار عددی ثابت انقباضی برابر با ۱/۷۳۰ خواهد بود. بنابراین در این مدل اگر مقادیر متغیرهای ورودی را به ترتیب ۵ = x_1 (میزان اسید استیک)، ۴ = x_2 (سولفید هیدروژن) و ۱ = x_3 (اسید لاکتیک) قرار دهیم نمره مربوط به مزه پنیر حدوداً ۲۰/۵۰۵ خواهد بود. شکل ۸.۳ نمودار مربوط به عدد فازی پاسخ مشاهده شده را نشان داده است.



شکل ۸.۳: تصویر پاسخ برآورد شده



شکل ۹.۳: مقایسه شهودی و استنباطی مدل‌های پایه و انقباضی در مقادیر مختلف k در مثال دوم: (آ) مقدار مرکز (ب) مقدار پهنای.

بررسی و ارزیابی عملکرد مدل‌های انقباضی نوع استاین بر اساس نمودار تیلور در این مثال، در شکل ۹.۳ بر اساس نمودار فوق میزان همبستگی و انحراف معیار تفاوت معنی داری

را نشان نمی‌دهند، البته در خصوص خطای استاندارد تولید شده، مدل‌های انقباضی نوع استاین خطای کمتری را نسبت به مدل مبنا نشان می‌دهد.

مثال سوم

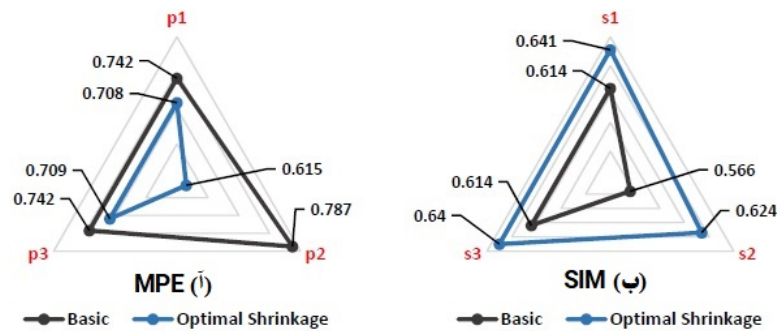
این مثال بر مبنای مدل تمام فازی انتخاب شده است، زیرا در مثال‌های قبل ورودی‌ها مقادیر غیر فازی بوده‌اند. درضمن داده‌های این مثال در بخش مثال‌های انگیزشی به عنوان سومین مجموعه داده در ۳.۲ قابل مشاهده است.

داده‌های این مثال در سال ۲۰۰۴ توسط نصرآبادی و همکاران در [۹۶] مدل‌بندی شده است و ما از آن به عنوان مدل مبنا بهره خواهیم برد. از این رو با بکارگیری روش انقباضی نوع استاین، نتایج ارزیابی مدل‌های انقباضی را در جدول ۱۰.۳ ارائه نموده‌ایم.

جدول ۱۰.۳: ارزیابی مدل‌های پایه و انقباضی تحت معیارهای شش گانه ارزیابی

معیارهای ارزیابی GOF						مدل
S_3	S_2	S_1	P_3	P_2	P_1	
۰/۶۱۴	۰/۵۶۶	۰/۶۱۴	۰/۷۴۲	۰/۷۸۷	۰/۷۴۲	مبنا [۹۶]
۰/۶۴۰	۰/۶۲۴	۰/۶۴۱	۰/۷۰۹	۰/۶۱۵	۰/۷۰۸	انقباضی
۰/۰۴۲	۰/۲۱۷	۰/۰۴۲	۰/۰۴۱	۰/۱۹۸	۰/۰۴۱	ثابت انقباضی بهینه (k)
(۰, ۰/۰۷۴]	(۰, ۰/۴۶۷]	(۰, ۰/۰۷۴]	(۰, ۰/۰۴۸]	(۰, ۰/۳۹۹]	(۰, ۰/۰۴۹]	محدوده بهینگی

بر این اساس طی نتایج حاصل از جدول ۱۰.۳ در بخش معیارهای بهینگی GOF، با کاهش مقادیر MPE و افزایش مقادیر SIM در مدل‌های انقباضی نسبت به مدل مبنا مواجه می‌شویم و این بیانگر برتری مدل به‌دست آمده با روش انقباضی نوع مثبت استاین نسبت به مدل مبناست. در واقع این روش بیان می‌کند که برآوردهای حاصله از هر مدل مبنا می‌تواند پایان کار برآزش نبوده و با این روش می‌توان به برآزش بهتری دست یافت. مقادیر k بهینه متناظر با هر کدام از معیارها و محدوده بهینگی آن‌ها نیز در این جدول ارائه شده است. لازم به ذکر است که اشتراک بازه‌های شش گانه بهینگی به صورت $(۰, ۰/۰۴۸]$ خواهد بود و بدین صورت مقدار عددی k برای تولید مدل IFS برابر با ۰/۰۴۸ به‌دست می‌آید. گزارش شهودی این برتری در شکل ۱۰.۳ قابل مشاهده است. همانطور که در شکل ۱.۲ مشاهده می‌شود، عملکرد مدل تولید شده توسط روش انقباضی در بخش (آ) که مثلث داخلی را ایجاد نموده، باعث کاهش معیارهای MPE در مقابل مدل مبنا است. همچنین این روش در بخش (ب) نیز با افزایش مقادیر معیارهای شباهت و ایجاد مثلث بزرگتر در این نمودار، شباهت بیشتر به مشاهدات واقعی را در برابر مدل مبنا القاء می‌کند.



شکل ۱۰.۳: مقایسه شهودی مدل پایه و انقباضی تحت معیارهای ارزیابی (آ) مقدار MPE (ب) مقدار SIM.

همچنین ضرایب تولید شده برای تمامی مدل‌های شش گانه، مدل مبنا و مدل IFS در قالب فرمول ریاضی رگرسیونی در جدول ۱۱.۳ آمده است.

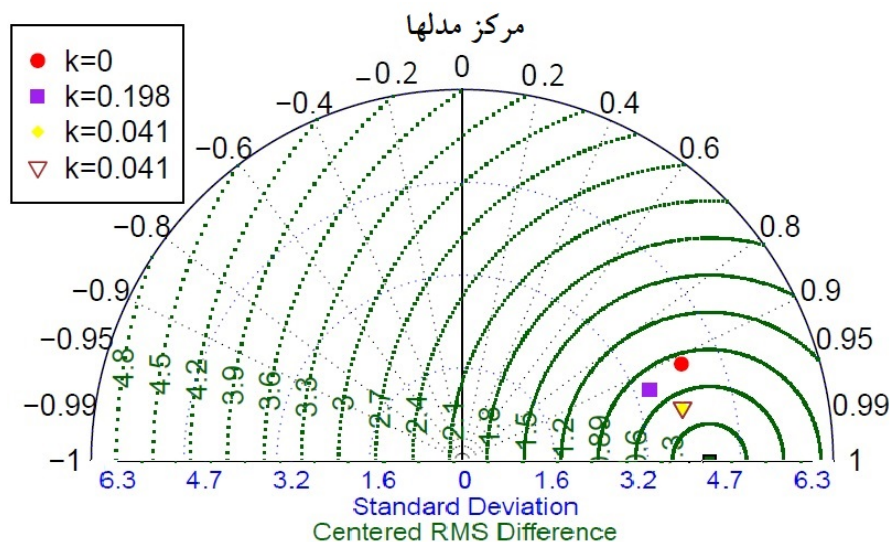
جدول ۱۱.۳: مدل‌های برازش شده توسط روش پایه، انقباضی نوع استاین و مدل IFS.

مدل	روش
$\tilde{y}_i = (3/580, 0/00)_T \oplus (0/550, 1/00)_T \otimes \tilde{x}_{i1}$	پایه [۹۶]
$\tilde{y}_i = (3/525, 0/00)_T \oplus (0/190, 0/802)_T \otimes \tilde{x}_{i1}$	۱
$\tilde{y}_i = (3/569, 0/00)_T \oplus (0/475, 0/959)_T \otimes \tilde{x}_{i1}$	۲
$\tilde{y}_i = (3/569, 0/00)_T \oplus (0/475, 0/959)_T \otimes \tilde{x}_{i1}$	۳
$\tilde{y}_i = (3/565, 0/00)_T \oplus (0/155, 0/783)_T \otimes \tilde{x}_{i1}$	۴
$\tilde{y}_i = (3/559, 0/00)_T \oplus (0/415, 0/926)_T \otimes \tilde{x}_{i1}$	۵
$\tilde{y}_i = (3/559, 0/00)_T \oplus (0/415, 0/926)_T \otimes \tilde{x}_{i1}$	۶
$\tilde{y}_i = (3/566, 0/00)_T \oplus (0/463, 0/952)_T \otimes \tilde{x}_{i1}$	IFS

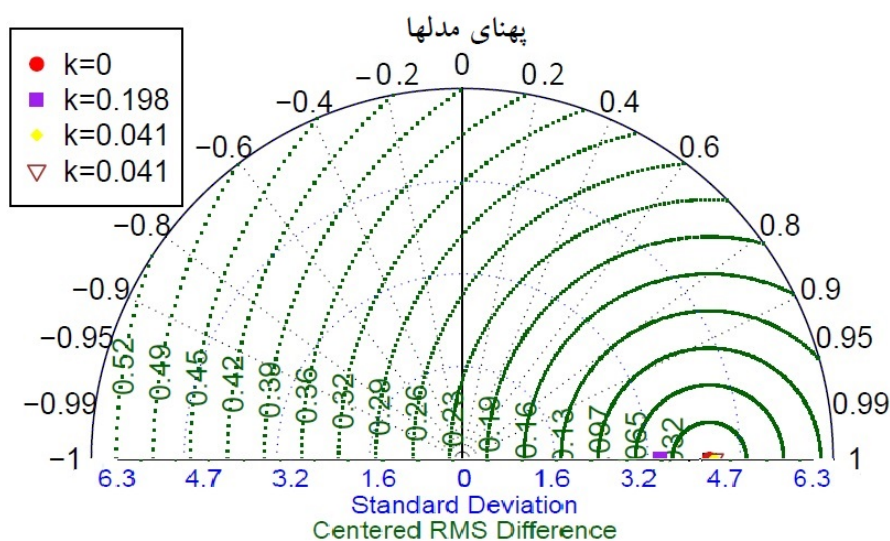
همچنین نمودار تیلور برای مقایسه مدل مبنا با مدل‌های انقباضی نوع استاین در خصوص تعامل با داده‌های واقعی نیز در شکل ۱۱.۳ قابل بررسی است.

بر اساس نمودار ۱۱.۳، در بخش (آ) مدل مبنا (دایره توپر) در قطاع بالاتری که عدد همبستگی پایین‌تری را بر روی محیط دایره نشان می‌دهد قرار دارد. لذا افزایش میزان همبستگی مشاهدات برآورد شده در مدل‌های انقباضی نوع استاین با مشاهدات واقعی که بر اساس مکان قرارگیری نمادهای مربوط به مدل‌های تولید شده، (مثلاً و مربع) توسط ثابت‌های انقباضی بهینه مختلف حاصل شده‌اند به وضوح مشاهده می‌شود. این عدد همبستگی در مدل‌های انقباضی حدوداً ۹۷٪ می‌باشد اما این شاخص برای مشاهدات برآورد شده در روش مبنا مقدار تقریبی ۹۱٪ را نشان می‌دهد. از سوی دیگر کاهش انحراف معیار و خطای استاندارد تولید شده توسط مدل‌های پیشنهادی نسبت به مدل مبنا به دلیل قرارگیری آنها در نیمه‌دایره‌های

پایین‌تر، حاکی از برتری مدل‌های انقباضی نوع استاین نسبت به مدل مبنا است. این کاهش در بخش (ب) نیز قابل مشاهده است. البته در خصوص مولفه میزان همبستگی اختلاف چندانی را نمی‌توان متصور شد.



(آ)



(ب)

شکل ۱۱.۳: مقایسه شهودی و استنباطی مدل‌های انقباضی در مقادیر مختلف k با مدل مبنا در مثال سوم: (آ) مقدار MPE (ب) مقدار SIM.

۴.۳ نتیجه‌گیری

در این فصل، ما یک استراتژی انقباضی را در مدل رگرسیون خطی فازی برای بهبود تخمین مدل ارائه کرده‌ایم. به عبارت دیگر، ما اجازه دادیم که برآوردهای فازی (به دست آمده از

هر روش فازی) به سمت مبدا منقبض شود و به طور خاص، ما از برآورد کننده انقباضی نوع استاین برای نقاط مرکزی و قسمت مثبت آن برای برآورد پهنای برآوردهای فازی استفاده کردیم. پس از آن ایده مدل IFS را برای یافتن مدل بهینه نهایی و یکتا در مواجهه با تمامی معیارهای ارزیابی GOF پیشنهاد کردیم. این مدل‌ها را در مجموعه‌ای از داده‌های شبیه‌سازی شده، بر روی ضرایب بدست آمده از مدل کمترین توان دوم که در سال ۱۹۹۸ مدل‌سازی شده بود اجرا کرده و برتری آن را نسبت به برخی از مدل‌های بهبود یافته اخیر با استفاده از معیارهای بهینگی GOF نشان دادیم. به بیانی دیگر، با این روش برای هرمدل تولید شده توسط هر روش می‌تواند پایان کار بهینگی مدل نباشد. ما همچنین از این روش برای تجزیه و تحلیل پاسخ‌ها در سه مثال واقعی با رویکردهای مختلف استفاده کردیم. علاوه بر این، برای اهداف استنباطی، ما نتیجه آزمایش فرضیه‌های فازی را برای ضرایب برآورد شده فازی مورد بررسی قرار دادیم.

با توجه به مثال‌های گویا، روش پیشنهادی به دلیل ارائه خطای پیش‌بینی میانگین کوچکتر (MPE) و معیارهای تشابه (SIM) بزرگتر به خوبی کار می‌کند. هزینه محاسباتی این روش نسبتاً کم است. یکی دیگر از مزایای روش پیشنهادی، بازه بهینگی ثابت انقباضی k است (نگاه کنید به رابطه ۵.۳). ممکن است یک متخصص تمایل داشته باشد که پهنایی با اندازه دلخواه داشته باشد. بازه‌های بهینه، به متخصص اجازه می‌دهد تا یک ثابت انقباضی را که هنوز هم برتری خود را در معیارهای مختلف ارزیابی حفظ کرده است، با یک پهنای مناسب انتخاب کند. رویکرد IFS به کاربران امکان می‌دهد تا با اشتراک‌گیری بین تمام بازه‌های بهینگی، مناسب‌ترین عدد ثابت انقباضی را که در تمامی معیارها دارای بهترین عملکرد است پیدا کنند.

برای تحقیقات بعدی، ابتدا اشاره می‌کنیم که روش پیشنهادی بسیار انعطاف‌پذیر است و به متخصص اجازه می‌دهد تا از روش منقبض شدن بر روی هر یک از مراکز، گسترش چپ یا راست با توجه به کاربرد، به طور جداگانه یا همزمان استفاده کند. از این رو، برای کارهای بیشتر، می‌توان این نتایج را برای مواردی که متخصص می‌خواهد برآوردها را برای رسیدن به حدس اولیه کوچک کند، برای اهداف عملی گسترش دهد. به علاوه، روش پیشنهادی در این فصل می‌تواند برای مدل‌های رگرسیونی فازی با جمله خطای فازی LR گسترش یابد (برای جزئیات بیشتر به [۳۱، ۷۷] مراجعه کنید).

فصل ۴

مدل‌های رگرسیونی جریمه شده فازی

۱.۴ مدل‌های جریمه‌ای

در این بخش در ادامه بحث بهینه‌سازی برآوردها، در دومین رویکرد به سراغ نحوه بکارگیری و عملکرد بهبود دهنده مدل‌های جریمه‌ای می‌پردازیم. در مدل‌بندی رگرسیونی ممکن است بین ستون‌های ماتریس طرح، همبستگی وجود داشته باشد. این ممکن است به واسطه ایجاد حالت‌های زیر باشد:

- افزایش تعداد متغیرهای پیشگو در ماتریس طرح.
- وجود رابطه‌ای خطی یا تقریباً خطی بین متغیرهای پیشگو.
- بیشتر بودن تعداد متغیرهای پیشگو از تعداد نمونه دریافتی که عملاً در آن ماتریس $(X^T X)^{-1}$ وجود ندارد.

یکی از نتایج آنها مشکل همخطی است که باعث افزایش واریانس پارامترهای برآورد شده می‌گردد و یعنی برآورد دقیقی از پارامترهای مدل را نخواهیم داشت. به گونه‌ای که ممکن محاسبه آنها با دو روش متفاوت نتایج کاملاً متفاوتی را به دست دهد. از این رو مدل‌های جریمه‌ای برای رفع این مشکل بسیار کارگشا خواهند بود. این گونه مدل‌ها در واقع با اضافه کردن جمله‌ای محدود کننده برای ضرایب مدل که تابعی از ضرایب متغیرهای ورودی است با نام جمله جریمه به ساختار عبارتی که فاصله بین پاسخ‌های مشاهده شده و برآورد شده را

محاسبه می‌کند، از افزایش واریانس برآورد گرهای جلوگیری نموده باعث کاهش میزان MPE مدل می‌شوند. البته عموماً این عمل آریبی پارامترهای برآورد شده را در پی دارد اما در توازن بین کاهش واریانس و افزایش آریبی، نسبت به MPE تولید شده توسط برآوردهای ناریب، نتیجه کار در مقادیر بهینه‌ای از جمله جریمه به سود کاهش MPE خواهد بود.

در بین مدل‌های جریمه‌ای با جملات محدود کنند، مدل‌های خاصی نیز وجود دارند که دارای خاصیت انتخاب متغیر هستند. انتخاب متغیر شیوه‌ای است که می‌تواند در ابعاد بالا اعمال شود بدین معنی که معمولاً تعداد متغیرهای توضیحی از تعداد مشاهدات نمونه بیشتر است یا همخطی معنی دار در بین متغیرهای ورودی وجود دارد و این یعنی متغیرهای پیش‌بینی کننده به طور خطی وابسته هستند و ماتریس $(X^T X)^{-1}$ یکتا نیست و یا اصلاً قابل محاسبه نیست.

شروع ساختار مدل‌های جریمه‌ای از هورل و کنارد [۶۷] آغاز گردید که رگرسیون رنج را معرفی کردند. آن‌ها با افزودن یک جمله جریمه به توان دوم خطای مدل رگرسیونی سعی کردند پارامترهای مدل رگرسیونی مربوط به رابطه متغیر پاسخ و متغیرهای ورودی با همبستگی بالا را تخمین بزنند. اگرچه روش رنج یک رهیافت برای مقابله با اثر همخطی است، اما در این روش امکان انتخاب متغیر وجود ندارد. در ادامه، روش کمترین انحراف عملگر انقباض و انتخاب کننده (لاسو (LASSO))، به‌عنوان ساختار جدیدی از مدل‌های جریمه‌ای توسط تیبشیرانی [۱۳۴] پیشنهاد شد که قادر به انتخاب متغیرهای کم اثر و حذف آن‌ها با صفر برآورد کردن ضرایب آن‌ها از مدل است. این روش، انتخاب متغیر را با استفاده از یک جمله جریمه قدرمطلق که جایگزین جمله جریمه توان دوم بکار رفته در روش رنج است، انجام می‌دهد (برای جزئیات بیشتر [۱۱۹] را ببینید).

بدنبال پیشرفت‌های اخیر در تکنیک‌های رنج و لاسو، رگرسیون بریج (Bridge) توسط [۵۵] که لاسو و رنج یکی از حالت‌های خاص آن هستند پیشنهاد گردید. در سال ۱۹۹۸ در [۵۶]، الگوریتمی برای حل برآورد ضرایب بریج ارائه گردید. همچنین در [۵۱]، روش جریمه قدرمطلق انحراف بریده شده (SCAD) پیشنهاد گردید که می‌تواند تضمین کننده راه حلی برای محاسبه ضرایب بزرگ اما غیر محدب که دارای تاثیر ناچیز، پیوسته و تقریباً ناریب هستند، باشد. در سال ۲۰۰۵ ژو و هستی [۱۵۳]، ترکیبی محدب از روش‌های جریمه‌ای رنج و لاسو را ارائه دادند که شبکه الاستیک (E-net) نامیده می‌شود. در واقع این ترکیب می‌تواند در موارد خاص به هر دو روش تبدیل شود (برای بررسی کامل این روش‌ها و سایر موارد در این موضوع، به [۱۴۷] مراجعه کنید).

از سوی دیگر در عمل ممکن است عدم دقت یا مبهم بودن در تعریف یا مشاهده متغیرهای پاسخ و یا توضیحی و یا عدم دقت در اندازه‌گیری پدیده‌ها به وجود آید. این ابهام ممکن است ناشی از بیان حوزه مقادیر متغیر به صورت تقریبی که ممکن است به واسطه عدم اشراف بر مقادیر متغیرها و یا کیفی بودن آن‌ها باشد. لذا، در چنین شرایطی استفاده از داده‌های فازی می‌تواند در بیان بهتر مقادیر متغیرها مفید باشد. لذا استفاده از مدل‌های جریمه‌ای

در رگرسیون فازی نیز دور از ذهن نیست.

همانطور که قبلاً اشاره شد، رگرسیون فازی برای اولین بار توسط تاناکا [۱۲۷] بر اساس یک مدل برنامه ریزی خطی، معروف به رگرسیون امکانی، معرفی شد. پس از آن، رویکردهای دیگری برای مدل‌سازی رگرسیون فازی تک متغیره (چند متغیره) ارائه گردید. این پیشنهادها شامل روش‌های امکانی، کمترین فاصله، ترکیبی و یا ابتکاری هستند. در برخی از روش‌ها، متغیرهای توضیحی یا اعداد فازی هستند مانند [۱۴۱، ۱۱۴، ۹۷، ۸۱، ۶۹، ۶۳، ۴۵، ۲۷، ۱۷] و یا اعداد غیر فازی هستند مانند [۱۵۰، ۱۲۲، ۱۰۴، ۸۱، ۵۹، ۴۳، ۴۰، ۳۴]. در این بین، دورسو و گاستالدی [۴۳] روشی را ارائه دادند که پاسخ را با ورودی‌های غیر فازی و خروجی فازی متقارن بر اساس فاصله اقلیدسی بین دو عدد فازی متقارن تولید می‌کند. از آنجا که گسترش یک عدد فازی می‌تواند ترکیبی خطی از درصدی از مرکز اعداد فازی باشد، آن‌ها در [۴۴]، پاسخ‌های تخمینی را با مدل‌های چند جمله‌ای برای دو نوع عدد فازی معرفی کردند. در [۳۶]، نویسندگان ابتدا یک مدل کلی فازی ارائه دادند. سپس آن‌ها موضوع انتخاب متغیر را در مدل‌های رگرسیون فازی با استفاده از ضریب تعیین تعدیل شده و متر پیشنهادی توسط یانگ و کو در [۱۴۳] معرفی کردند. استفاده از این ایده در رگرسیون خطی وزن‌دار شده فازی و استوار به ترتیب در [۴۵] و [۴۶] مشاهده می‌شود.

بررسی مدل‌های جریمه شده در رگرسیون فازی نیز انجام گردید. هونگ و همکاران [۷۰] برآورد ریج را با استفاده از متغیرهای ورودی غیر فازی و یک عدد فازی نرمال به عنوان پاسخ پیشنهاد دادند. فرنوش و همکاران [۵۲] نیز یک ابزار انتخاب متغیر به نام برآوردگر نوع ریج برای مدل‌های رگرسیون فازی با استفاده از متغیرهای ورودی و خروجی فازی پیشنهاد دادند. در [۱۱۰]، ربیعی و همکاران در یک محیط فازی، یک رگرسیون ریج دو پارامتری را مورد بررسی قرار دادند. همچنین اخیراً، حسامیان و اکبری [۶۵] نیز مدل رگرسیون LASSO فازی را ارائه دادند. بنابراین با در نظر گرفتن اهمیت انتخاب متغیر در مدل رگرسیون فازی، در این بخش مدل رگرسیون خطی فازی جریمه شده برای ورودی‌های غیر فازی و خروجی LR -فازی را ارائه خواهیم نمود.

در مقایسه با تحقیقات ذکر شده مزیت مطالعه ما عمومیت مدل جریمه‌ای است که شامل ریج و لاسو به عنوان موارد خاص است. علاوه بر این، روش پیشنهادی ما صورت بسته‌ای از مدل جریمه‌ای را در برآورد پارامترهای مدل ارائه می‌دهد. در این مسیر، با استفاده از فاصله پیشنهادی توسط یانگ و کو [۱۴۳] و اضافه کردن تابعی کلی از ضرایب به عنوان جمله جریمه و با استفاده از رویکرد [۴۲] صورت بسته‌ای را برای برآورد ضرایب و همزمان انتخاب متغیر تولید می‌نماییم. سپس در دو بخش شامل مثال‌های شبیه‌سازی شده در سناریوهای مختلف همبستگی و مثال‌های واقعی با استفاده از معیارهای بهینگی GOF مدل پیشنهادی خود را در حالت‌های مختلفی از تابع جریمه با مدل‌های دیگر مقایسه خواهیم کرد. همچنین دقت مدل پیشنهادی را در داده‌های واقعی توسط بوت استرپ ناپارامتری ارزیابی می‌کنیم.

۲.۴ مدل رگرسیون جریمه شده فازی

در مدل‌های رگرسیون فازی (مشابه رگرسیون کلاسیک) ممکن است موارد زیر برای متغیرهای توضیحی رخ دهد تا در خصوص بکارگیری از مدل‌های جریمه‌ای اقدام نماییم:

الف) برخی از متغیرهای ورودی ممکن است کمترین تأثیر را بر روی مدل داشته باشند. بنابراین، می‌توانیم آن‌ها را به شدت غیر موثر بنامیم.

ب) ممکن است بین متغیرهای ورودی همبستگی معنی داری وجود داشته باشد (که مطمئناً با افزایش تعداد متغیرهای ورودی این مسئله بیشتر امکان‌پذیر است).

بنابراین، استفاده از روش‌های جریمه شده نه تنها می‌تواند با صفر برآورد کردن ضرایب فازی آن‌ها برخی از متغیرهای ورودی غیر موثر را حذف کند، بلکه می‌تواند معیارهای بهینگی مدل (GOF) را در مقایسه با سایر مدل‌ها بهبود بخشد. مدل رگرسیون فازی را به صورت زیر در نظر بگیرید:

$$\tilde{y} = \mathbf{X} \otimes \tilde{\beta} \oplus \tilde{\varepsilon} \quad (1.4)$$

که در آن $\tilde{y} = (y, s, r)_{LR}$ شامل بردارهای r, s, y با ابعاد $1 \times (n+1)$ و به ترتیب نمایانگر مراکز، پهناهای راست و چپ اعداد فازی هستند، بردارهایی با بعد $\tilde{\beta} = (\beta, \nu_1, \nu_2)_{LR}$ و به عنوان ضریب مدل رگرسیونی شامل بردارهای $1 \times (p+1)$ بعدی β, ν_1, ν_2 که به ترتیب نمایانگر مراکز، پهناهای راست و چپ اعداد فازی هستند. همچنین واحد $\tilde{\varepsilon} = (\varepsilon_y, \varepsilon_s, \varepsilon_r)$ نیز به عنوان بردار واحد $1 \times n$ بعدی خطای مدل که در آن $\varepsilon_r, \varepsilon_s, \varepsilon_y$ به ترتیب نمایانگر مراکز، پهناهای راست و چپ اعداد فازی هستند. ماتریس $\mathbf{X}$ هم به عنوان ماتریس ورودی‌ها، دارای ابعاد $(p+1) \times n$ و شامل یک ستون با مقدار ۱ و p ستون برای مقادیر متغیرهای توضیحی است. بنابراین با در نظر گرفتن $\hat{\tilde{y}} = (\hat{y}, \hat{s}, \hat{r})_{LR}$ به عنوان پاسخ برآورد شده می‌توان نوشت:

$$\begin{cases} y = \hat{y} + \varepsilon_y, & \hat{y} = \mathbf{X}h \\ s = \hat{s} + \varepsilon_s, & \hat{s} = \mathbf{X}ha + 1b \\ r = \hat{r} + \varepsilon_r, & \hat{r} = \mathbf{X}hc + 1f \end{cases} \quad (2.4)$$

که در آن $\hat{y}, \hat{s}$ و ۱ همگی بردارهای $1 \times n$ ، h یک بردار $1 \times p$ و پارامترهای a, b, c و f مقادیری حقیقی هستند.

توجه داشته باشید که به دنبال رویکرد دورسو [۴۲]، مدل رگرسیون خطی فازی پیشنهادی متشکل از مدل‌سازی همزمان مرکز متغیر خروجی فازی LR با استفاده از یک مدل رگرسیون چندگانه بر روی متغیرهای ورودی غیر فازی، پهناهای چپ و راست متغیر خروجی از طریق

دو رگرسیون خطی چندگانه و وابسته به مراکز برآورده شده است. اکنون، فاصله جریمه شده عبارتست از:

$$\psi(\mathbf{h}, a, b, c, f) = D^2(\tilde{\mathbf{y}}, \hat{\mathbf{y}}) + \lambda \sum_{j=1}^p \phi(|h_j|), \quad \lambda > 0 \quad (3.4)$$

که در آن $D^2(\tilde{\mathbf{y}}, \hat{\mathbf{y}})$ فاصله تعریف شده در بخش ۴.۱ بین پاسخ برآورد شده و مشاهده شده است. همچنین $\phi(\cdot)$ نماد مربوط به هر تابع بولر اندازه‌پذیر است و در واقع این بخش مربوط به ساختار جریمه‌ای مدل می‌باشد. پارامتر λ نیز به عنوان تنظیم کننده میزان انقباض در جمله جریمه است. $\psi(\mathbf{h}, a, b, c, f)$ نیز به عنوان فاصله بین متغیرهای پاسخ فاز مشاهده شده و برآورد شده است و بر اساس رابطه (۴.۱) تشریح می‌شود. در این بخش برای بدست آوردن برآورد پارامترهای مدل لم زیر را بیان و اثبات می‌کنیم.

لم ۱.۲.۴. در رگرسیون فاز با متغیرهای ورودی غیر فاز و خروجی فاز LR، برآورد پارامترهای مدل رگرسیونی جریمه شده

$$\psi(\mathbf{h}, a, b, c, f) = D^2(\tilde{\mathbf{y}}, \hat{\mathbf{y}}) + \lambda \sum_{j=1}^p \phi(|h_j|), \quad \lambda > 0$$

با استفاده از رویکرد پیشنهادی در [۴۲] و متر پیشنهادی در [۱۴۳] عبارتست از:

$$\begin{aligned} \hat{\mathbf{h}} &= ((\mathbf{I} + G_1)\mathbf{X}^T\mathbf{X} + \lambda\mathbf{Q}_0)^{-1}((\mathbf{I} + G_7)\mathbf{X}^T\mathbf{y} + G_7\mathbf{X}^T + G_4\mathbf{X}^T\mathbf{1}) \\ \hat{a} &= (k_1\hat{\mathbf{h}}^T\mathbf{X}^T\mathbf{X}\hat{\mathbf{h}})^{-1}(\hat{\mathbf{h}}^T\mathbf{X}^T\mathbf{X}\hat{\mathbf{h}} - \mathbf{y}^T\mathbf{X}\hat{\mathbf{h}} + k_1\mathbf{s}^T\mathbf{X}\hat{\mathbf{h}} - k_1\mathbf{1}^T\mathbf{X}\hat{\mathbf{h}}) \\ \hat{b} &= \frac{1}{k_1n}(k_1\mathbf{1}^T\mathbf{s} - \mathbf{1}^T\mathbf{y} + \mathbf{1}^T\mathbf{X}\hat{\mathbf{h}} - k_1a\mathbf{1}^T\mathbf{X}\hat{\mathbf{h}}) \\ \hat{c} &= (k_7\hat{\mathbf{h}}^T\mathbf{X}^T\mathbf{X}\hat{\mathbf{h}})^{-1}(-\hat{\mathbf{h}}^T\mathbf{X}^T\mathbf{X}\hat{\mathbf{h}} + \mathbf{y}^T\mathbf{X}\hat{\mathbf{h}} + k_1\mathbf{r}^T\mathbf{X}\hat{\mathbf{h}} - k_1f\mathbf{1}^T\mathbf{X}\hat{\mathbf{h}}) \\ \hat{f} &= \frac{1}{k_7n}(k_7\mathbf{1}^T\mathbf{r} + \mathbf{1}^T\mathbf{y} - \mathbf{1}^T\mathbf{X}\hat{\mathbf{h}} - k_7c\mathbf{1}^T\mathbf{X}\hat{\mathbf{h}}) \end{aligned}$$

که در آن $\mathbf{Q}_0 = \text{Diag} \left\{ \frac{\phi'(h_1^{(c)})}{h_1^{(c)}}, \dots, \frac{\phi'(h_p^{(c)})}{h_p^{(c)}} \right\}$ و همچنین

$$\begin{aligned} G_1 &= -\frac{1}{\lambda}k_1a + \frac{1}{\lambda}k_1^2a^2 + \frac{1}{\lambda}k_7c + \frac{1}{\lambda}k_7^2c^2, & G_7 &= -\frac{1}{\lambda}k_1a + \frac{1}{\lambda}k_7c \\ G_7 &= (\frac{1}{\lambda}k_1^2a - \frac{1}{\lambda}k_1)\mathbf{s} + (\frac{1}{\lambda}k_7^2c + \frac{1}{\lambda}k_7)\mathbf{r}, & G_4 &= -\frac{1}{\lambda}k_1b + \frac{1}{\lambda}k_1^2ab - \frac{1}{\lambda}k_7f - \frac{1}{\lambda}k_7^2cf \end{aligned}$$

هستند.

برهان: همانطور که در فرضیات لم ۱.۲.۴ آمده است اکنون با استفاده از رابطه فاصله (۴.۱)، معادله (۳.۴) به صورت زیر بازنویسی می‌شود:

$$\begin{aligned}
 \psi(\mathbf{h}, a, b, c, f) = & \frac{1}{\Psi} [(\mathbf{y} - \mathbf{Xh})^\top (\mathbf{y} - \mathbf{Xh}) \\
 & + (\mathbf{y} - \mathbf{Xh} - k_1 \mathbf{s} + k_1 a \mathbf{Xh} + k_1 b \mathbf{1})^\top (\mathbf{y} - \mathbf{Xh} - k_1 \mathbf{s} + k_1 a \mathbf{Xh} + k_1 b \mathbf{1}) \\
 & + (\mathbf{y} - \mathbf{Xh} + k_2 \mathbf{r} - k_2 c \mathbf{Xh} - k_2 f \mathbf{1})^\top (\mathbf{y} - \mathbf{Xh} + k_2 \mathbf{r} - k_2 c \mathbf{Xh} - k_2 f \mathbf{1})] \\
 & + \lambda \sum_{j=1}^p \phi(|h_j|)
 \end{aligned} \quad (۴.۴)$$

اکنون با مشتق‌گیری نسبت به $\mathbf{h}$ و حل آن داریم:

$$\begin{aligned}
 \frac{\partial \psi(\mathbf{h}, a, b, c, f)}{\partial \mathbf{h}} = & \frac{\partial}{\partial \mathbf{h}} \left[\frac{1}{\Psi} [\mathbf{y}^\top \mathbf{y} - \mathbf{y}^\top \mathbf{Xh} - \mathbf{h}^\top \mathbf{X}^\top \mathbf{y} + \mathbf{h}^\top \mathbf{X}^\top \mathbf{Xh} + \mathbf{y}^\top \mathbf{y} \right. \\
 & - k_1 \mathbf{y}^\top \mathbf{s} - \mathbf{y}^\top \mathbf{Xh} - k_1 a \mathbf{y}^\top \mathbf{Xh} + k_1 b \mathbf{y}^\top \mathbf{1} - k_1 \mathbf{s}^\top \mathbf{y} + k_1 \mathbf{s}^\top \mathbf{s} \\
 & + k_1 \mathbf{s}^\top \mathbf{Xh} - k_1 a \mathbf{s}^\top \mathbf{Xh} - k_1 b \mathbf{s}^\top \mathbf{1} - \mathbf{h}^\top \mathbf{X}^\top \mathbf{y} + k_1 \mathbf{h}^\top \mathbf{X}^\top \mathbf{s} \\
 & + \mathbf{h}^\top \mathbf{X}^\top \mathbf{Xh} - k_1 a \mathbf{h}^\top \mathbf{X}^\top \mathbf{Xh} - k_1 b \mathbf{h}^\top \mathbf{X}^\top \mathbf{1} + k_1 a \mathbf{h}^\top \mathbf{X}^\top \mathbf{y} \\
 & - k_1 a \mathbf{h}^\top \mathbf{X}^\top \mathbf{s} - k_1 a \mathbf{h}^\top \mathbf{X}^\top \mathbf{Xh} + k_1 a^2 \mathbf{h}^\top \mathbf{X}^\top \mathbf{Xh} + k_1 a b \mathbf{h}^\top \mathbf{X}^\top \mathbf{1} \\
 & + k_1 b \mathbf{1}^\top \mathbf{y} - k_1 b \mathbf{1}^\top \mathbf{s} - k_1 b \mathbf{1}^\top \mathbf{Xh} + k_1 a b \mathbf{1}^\top \mathbf{Xh} + k_1 b^2 \mathbf{1}^\top \mathbf{1} \\
 & + \mathbf{y}^\top \mathbf{y} + k_2 \mathbf{y}^\top \mathbf{r} - \mathbf{y}^\top \mathbf{Xh} - k_2 c \mathbf{y}^\top \mathbf{Xh} - k_2 f \mathbf{y}^\top \mathbf{1} + k_2 \mathbf{r}^\top \mathbf{y} \\
 & + k_2 \mathbf{r}^\top \mathbf{r} - k_2 \mathbf{r}^\top \mathbf{Xh} - k_2 c \mathbf{r}^\top \mathbf{Xh} + k_2 f \mathbf{r}^\top \mathbf{1} - \mathbf{h}^\top \mathbf{X}^\top \mathbf{y} \\
 & - k_2 \mathbf{h}^\top \mathbf{X}^\top \mathbf{r} + \mathbf{h}^\top \mathbf{X}^\top \mathbf{Xh} + k_2 c \mathbf{h}^\top \mathbf{X}^\top \mathbf{Xh} + k_2 f \mathbf{h}^\top \mathbf{X}^\top \mathbf{1} \\
 & - k_2 c \mathbf{h}^\top \mathbf{X}^\top \mathbf{y} - k_2 c \mathbf{h}^\top \mathbf{X}^\top \mathbf{r} + k_2 c \mathbf{h}^\top \mathbf{X}^\top \mathbf{Xh} + k_2 c^2 \mathbf{h}^\top \mathbf{X}^\top \mathbf{Xh} \\
 & + k_2 c f \mathbf{h}^\top \mathbf{X}^\top \mathbf{1} - k_2 f \mathbf{1}^\top \mathbf{y} - k_2 f \mathbf{1}^\top \mathbf{r} + k_2 f \mathbf{1}^\top \mathbf{Xh} + k_2 c f \mathbf{1}^\top \mathbf{Xh} \\
 & \left. + k_2 f^2 \mathbf{1}^\top \mathbf{1}] + \lambda \sum_{j=1}^p \phi(|h_j|) \right] = 0
 \end{aligned} \quad (۵.۴)$$

اکنون با ساده کردن عبارات فوق خواهیم داشت:

$$\begin{aligned}
 \frac{\partial \psi(\mathbf{h}, a, b, c, f)}{\partial \mathbf{h}} = & \frac{1}{\Psi} [-\mathbf{X}^\top \mathbf{y} + \mathbf{X}^\top \mathbf{Xh} + \mathbf{1} k_1 a \mathbf{X}^\top \mathbf{y} + \mathbf{1} k_1 \mathbf{X}^\top \mathbf{s} - \mathbf{1} k_1 a \mathbf{X}^\top \mathbf{s} \\
 & - \mathbf{1} k_1 a \mathbf{X}^\top \mathbf{Xh} - \mathbf{1} k_1 b \mathbf{X}^\top \mathbf{1} + \mathbf{1} k_1 a^2 \mathbf{X}^\top \mathbf{Xh} + \mathbf{1} k_1 a b \mathbf{X}^\top \mathbf{1} \\
 & - \mathbf{1} k_2 c \mathbf{X}^\top \mathbf{y} - \mathbf{1} k_2 \mathbf{X}^\top \mathbf{r} - \mathbf{1} k_2 c \mathbf{X}^\top \mathbf{r} - \mathbf{1} k_2 c \mathbf{X}^\top \mathbf{Xh} + \mathbf{1} k_2 f \mathbf{X}^\top \mathbf{1} \\
 & + \mathbf{1} k_2 c^2 \mathbf{X}^\top \mathbf{Xh} + \mathbf{1} k_2 c f \mathbf{X}^\top \mathbf{1}] + \lambda \sum_{j=1}^p (\phi(|h_j|))' = 0
 \end{aligned} \quad (۶.۴)$$

که در آن $(\phi(|h_j|))'$ مشتق تابع ϕ نسبت به h_j است.

پس از مرتب کردن عبارات (۶.۴) نسبت به $\mathbf{h}$ داریم:

$$\begin{aligned} (1 - \frac{1}{\beta} k_1 a + \frac{1}{\beta} k_1^* a^* + \frac{1}{\beta} k_2 c + \frac{1}{\beta} k_2^* c^*) \mathbf{X}^T \mathbf{X} \mathbf{h} + \lambda \sum_{j=1}^p (\phi(|h_j|))' = \\ (1 - \frac{1}{\beta} k_1 a + \frac{1}{\beta} k_2 c) \mathbf{X}^T \mathbf{y} + ((\frac{1}{\beta} k_1^* a - \frac{1}{\beta} k_1) \mathbf{s} + (\frac{1}{\beta} k_2^* c + \frac{1}{\beta} k_2) \mathbf{r}) \mathbf{X}^T \\ + (\frac{1}{\beta} k_1 b - \frac{1}{\beta} k_1^* ab - \frac{1}{\beta} k_2 f - \frac{1}{\beta} k_2^* cf) \mathbf{X}^T \mathbf{1} \end{aligned} \quad (7.4)$$

اکنون برای بدست آوردن یک صورت بسته شده بر اساس $\mathbf{h}$ ، فرض کنید مقدار اولیه $\mathbf{h}^{(0)}$ به ما داده شود که نزدیک به مقدار کمینه رابطه (۴.۴) باشد. اگر $h_j^{(0)}$ بسیار نزدیک به صفر باشد آنگاه $h_j = 0$ ، $j = 1, \dots, p$ فرض می‌شود (برای اطلاعات بیشتر به [۵۱] مراجعه کنید). بنابراین، تابع جریمه را می‌توان با یک تابع درجه دوم به صورت محلی تقریب زد

$$(\phi(|h_j|))' = \phi'(h_j) \text{sgn}(h_j) \approx \left\{ \frac{\phi'(h_j^{(0)})}{h_j^{(0)}} \right\} h_j, \quad j = 1, \dots, p \quad (8.4)$$

که در آن $\text{sgn}(z)$ تابع علامت است و برای مقادیر $z > 0$ عدد یک و برای مقادیر $z < 0$ عدد صفر را دریافت می‌کند.

اکنون با جایگذاری (۸.۴) در (۷.۴) صورت بسته ی برآورد $\mathbf{h}$ به صورت زیر حاصل می‌شود:

$$\hat{\mathbf{h}} = ((1 + G_1) \mathbf{X}^T \mathbf{X} + \lambda \mathbf{Q}_0)^{-1} ((1 + G_2) \mathbf{X}^T \mathbf{y} + G_3 \mathbf{X}^T + G_4 \mathbf{X}^T \mathbf{1}) \quad (9.4)$$

حال با داشتن $\hat{\mathbf{h}}$ ، برآورد پارامترهای a, b, c ، و f بر اساس مشتق گیری نسبت به هر کدام از آنها و جایگذاری $\hat{\mathbf{h}}$ در روابط مربوط به آنها عبارتست از:

$$\hat{a} = (k_1 \hat{\mathbf{h}}^T \mathbf{X}^T \mathbf{X} \hat{\mathbf{h}})^{-1} (\hat{\mathbf{h}}^T \mathbf{X}^T \mathbf{X} \hat{\mathbf{h}} - \mathbf{y}^T \mathbf{X} \hat{\mathbf{h}} + k_1 \mathbf{s}^T \mathbf{X} \hat{\mathbf{h}} - k_1 b \mathbf{1}^T \mathbf{X} \hat{\mathbf{h}}) \quad (10.4)$$

$$\hat{b} = \frac{1}{k_1 n} (k_1 \mathbf{1}^T \mathbf{s} - \mathbf{1}^T \mathbf{y} + \mathbf{1}^T \mathbf{X} \hat{\mathbf{h}} - k_1 a \mathbf{1}^T \mathbf{X} \hat{\mathbf{h}}) \quad (11.4)$$

$$\hat{c} = (k_2 \hat{\mathbf{h}}^T \mathbf{X}^T \mathbf{X} \hat{\mathbf{h}})^{-1} (-\hat{\mathbf{h}}^T \mathbf{X}^T \mathbf{X} \hat{\mathbf{h}} + \mathbf{y}^T \mathbf{X} \hat{\mathbf{h}} + k_1 \mathbf{r}^T \mathbf{X} \hat{\mathbf{h}} - k_1 f \mathbf{1}^T \mathbf{X} \hat{\mathbf{h}}) \quad (12.4)$$

$$\hat{f} = \frac{1}{k_2 n} (k_2 \mathbf{1}^T \mathbf{r} + \mathbf{1}^T \mathbf{y} - \mathbf{1}^T \mathbf{X} \hat{\mathbf{h}} - k_2 c \mathbf{1}^T \mathbf{X} \hat{\mathbf{h}}) \quad (13.4)$$

□

حال برآوردگر جریمه‌ای در حالات خاص تابع جریمه به صورت زیر بررسی می‌کنیم:

بریج: در این حالت $\phi(|h_j|) = |h_j|^t$ ، برای $t > 0$ و $j = 1, \dots, p$ را در نظر می‌گیریم. آنگاه داریم: $(\phi(|h_j|))' = t|h_j|^{t-1} \text{sgn}(h_j)$.

لاسو: در اینجا $\phi(|h_j|) = |h_j|$ را در نظر می‌گیریم که در واقع معادل قرار دادن $t = 1$ در توان تابع جریمه بریج است. پس داریم: $(\phi(|h_j|))' = t|h_j| \text{sgn}(h_j)$.

ریج: اکنون در نظر می‌گیریم: $\phi(|h_j|) = h_j^2$ ، که معادل قرار دادن $t = 2$ بجای توان تابع جریمه بریج است. پس داریم: $(\phi(|h_j|))' = 2h_j$.

حال با قرار دادن هرکدام از توابع بالا بجای $\hat{h}$ به عنوان برآورد h به ترتیب هر یک عبارتست از:

$$\hat{h} = ((1 + G_1)X^T X + \lambda Q_o^B)^{-1} ((1 + G_2)X^T y + G_3 X^T + G_4 X^T 1) \quad (14.4)$$

$$\hat{h} = ((1 + G_1)X^T X + \lambda Q_o^L)^{-1} ((1 + G_2)X^T y + G_3 X^T + G_4 X^T 1) \quad (15.4)$$

$$\hat{h} = ((1 + G_1)X^T X + \lambda I)^{-1} ((1 + G_2)X^T y + G_3 X^T + G_4 X^T 1) \quad (16.4)$$

که در آن $Q_o^L = \text{Diag} \left\{ \frac{|h_1^{(o)}|}{h_1^{(o)}}, \dots, \frac{|h_p^{(o)}|}{h_p^{(o)}} \right\}$ و $Q_o^B = \text{Diag} \left\{ \frac{|h_1^{(o)}|^{t-1}}{h_1^{(o)}}, \dots, \frac{|h_p^{(o)}|^{t-1}}{h_p^{(o)}} \right\}$ است.

تبصره ۱.۲.۴. توجه داشته باشید برآوردهای بدست آمده در ۱۰.۴ تا ۱۳.۴ به طور خودکار عدم منفی بودن پهناهای تخمین زده شده در مدل ۴.۴ را تضمین نمی‌کنند. برای رفع این مشکل از آنجایی که پهناهای عدد فازی می‌بایست نامنفی باشد برای هر $\tilde{\beta}_i$ شرط نامنفی بودن را به صورت $\nu_i = \max(\circ, \nu_i)$ در نظر خواهیم گرفت. این رویکرد در [۴۲] پیشنهاد شده است.

تبصره ۲.۲.۴. بر اساس روابط بالا در صورتی که $a = c = \circ$ ، $b = f = \circ$ و $k_1 = k_2$ ، یعنی با حذف پهناها به حالت غیر فازی (حالت تولید شده در آمار کلاسیک) خواهیم رسید.

تبصره ۳.۲.۴. برای اعداد فازی مثلثی و دوزنقه‌ای کفایت قرار دهیم $k_1 = k_2 = \frac{1}{2}$.

تبصره ۴.۲.۴. پارامتر تنظیم λ نقشی تعیین کننده در این روش دارد و انتخاب آن باید با دقت انجام شود. وقتی λ بزرگتر می‌شود، با توجه به انتخاب مدل، پارامترهای بیشتری روی صفر تنظیم می‌شوند از این رو، این پارامتر را به گونه‌ای انتخاب می‌کنیم که میانگین خطای پیش‌بینی (MPE) کمتری داشته باشیم. در واقع مانند تنظیم پیچ رادیو به دنبال λ ای هستیم که در آن کمترین MPE (به عنوان معیار بهینگی مدل) را داشته باشد.

۳.۴ معیارهای بهینگی مدل

برای بررسی و مقایسه مدل پیشنهادی خود با دیگر مدل‌ها، در خصوص نمایش برتری آن، از دو معیار بهینگی مدل (GOF) استفاده می‌شود.

فرض کنید $\tilde{y} = (y, s, r)_{LR}$ و $\hat{\tilde{y}} = (\hat{y}, \hat{s}, \hat{r})_{LR}$ به ترتیب پاسخ‌های مشاهده شده و برآورد شده باشند. اکنون بر اساس معادله (۱.۴) داریم

$$\hat{\tilde{y}} = X \otimes \hat{\tilde{\beta}} \implies X^T \hat{\tilde{y}} = X^T X \otimes \hat{\tilde{\beta}} \implies \hat{\tilde{\beta}} = (X^T X)^{-1} X^T \hat{\tilde{y}} \quad (17.4)$$

اکنون، برای ورود به روش انتخاب متغیر، ابتدا مقادیر مختلف $\hat{\tilde{y}}$ را با تغییر مقدار $\lambda > \circ$ به عنوان پارامتر تنظیم کننده به دست می‌آوریم. برای مقایسه روشهای جریمه شده با دیگر

روش‌ها، به معیارهایی نیاز داریم. از آنجا که هدف اصلی رگرسیون خطی فازی (FLR) پیش‌بینی است، ما میانگین خطای پیش‌بینی (MPE) را به عنوان معیاری برای اختلاف بین مقادیر برآورد شده و مشاهده شده متغیر پاسخ به صورت زیر محاسبه می‌کنیم.

الف) میانگین خطای پیش‌بینی (MPE):

$$P = \text{MPE}(\tilde{y}, \hat{y}) = \frac{1}{n} \sum_{i=1}^n D_{\gamma}(\tilde{y}_i, \hat{y}_i)$$

که $D_{\gamma}(\tilde{y}_i, \hat{y}_i)$ ریشه دوم مقدار فاصله (۶.آ) بین $\tilde{y}_i$ و $\hat{y}_i$ است. همچنین اعتباریابی متقابل^۱ (CV) عبارتست از:

$$CV = \frac{1}{n} \sum_{i=1}^n D_{\gamma}(\tilde{y}_i, \hat{y}_i^*)$$

و در آن $\tilde{y}_i$ پاسخ‌های مشاهده شده و $\hat{y}_i^*$ پاسخ برآورد شده‌ای است که در آن i – امین نمونه از مجموعه داده‌ها کنار گذاشته شده است.

ب) معیار شباهت (SIM): این معیار نیز به صورت زیر و براساس فاصله ارائه شده در ۶.آ تعریف می‌گردد:

$$S = \text{SIM}(\tilde{y}, \hat{y}) = \frac{1}{n} \sum_{i=1}^n \frac{1}{1 + D_{\gamma}(\tilde{y}_i, \hat{y}_i)}$$

۱.۳.۴ معیار نایقینی

برای بررسی عدم دقت به عنوان یک معیار نایقینی، به دلیل عدم اطلاع کافی در مورد روند تولید داده، از یک رویکرد بوت استرپ غیرپارامتری استفاده می‌کنیم بدین معنی که انحراف معیار مرکز و پهنای چپ و راست هر ضریب مدل رگرسیونی در نمونه‌های بوت استرپ را محاسبه می‌کنیم (این ساز و کار در [۳۶، ۴۶] اجرا شده است).
فرآیند اجرایی این معیار در الگوریتم ۴ آمده است.

^۱Cross Validation

الگوریتم ۴ الگوریتم محاسباتی برای روش بوت استرپ

- گام ۱:** تعداد B نمونه بوت استرپ را در نظر بگیرید.
- گام ۲:** یک نمونه تصادفی با جایگذاری به اندازه B از خطای برآورد شده $\hat{\epsilon}^{(d)} = \tilde{y}^{(d)} - \hat{y}^{(d)}$ ، $d = 1, \dots, B$ تولید کنید (در اینجا شاخص d بیانگر، d امین نمونه بوت استرپ است).
- گام ۳:** پاسخ‌های مشاهده شده جدید را با استفاده از نمونه‌های گام دوم به صورت زیر به دست آورید.

$$\hat{\mathbf{y}}^{\bullet(d)} = \hat{\epsilon}_i^{(d)} + \hat{\mathbf{y}}_i^{(d)}, d = 1, \dots, B$$

- گام ۴:** $\tilde{\beta}^{(d)}$ ها را با استفاده از پاسخ‌های مشاهده شده $\hat{\mathbf{y}}^{\bullet(d)} = (\hat{y}_1^{\bullet(d)}, \dots, \hat{y}_n^{\bullet(d)})^\top$ برای به دست آوردن $\{\tilde{\beta}^{(1)}, \tilde{\beta}^{(2)}, \dots, \tilde{\beta}^{(B)}\}$ برآورد کنید. در پایان برآورد بوت استرپ به صورت زیر تهیه می‌شود.

$$\hat{\beta}^B = \frac{1}{B} \sum_{d=1}^B \tilde{\beta}^{(d)}$$

۲.۳.۴ فاصله اطمینان t - بوت استرپی

بوت استرپ غیرپارامتری امکان تخمین توزیع نمونه‌ای ضرایب رگرسیون را به صورت تجربی و بدون پیش فرض در مورد توزیع جامعه فراهم می‌کند. فاصله اطمینان بوت استرپی^۲ (CI) برای هر ضریب رگرسیون محاسبه می‌شود. افرون و تیبشیرانی [۴۸] مدل‌های مختلفی از ساختارهای فاصله اطمینان بوت استرپی مانند فاصله اطمینان چندکی و فاصله اطمینان t - بوت استرپی را ارائه نموده اند. بدیهی است، اگر CI حاوی صفر باشد، فرضیه صفر، یعنی صفر بودن ضریب رگرسیون مربوطه قابل رد نیست. در اینجا، ما فاصله اطمینان t - بوت استرپی [۴۸] را به طور خلاصه در الگوریتم ۵ شرح می‌دهیم. همچنین این ساختار توسط [۸۷] استفاده شده است.

^۲Confidence interval

الگوریتم ۵ فاصله اطمینان t - بوت استری

گام ۱: تعداد B نمونه بوت استری در نظر بگیرید.

گام ۲: هر بخش از $\hat{\beta}_j^{(d)} = (\hat{m}_{\beta_j^{(d)}}, \hat{l}_{\beta_j^{(d)}}, \hat{r}_{\beta_j^{(d)}})_{LR}$ را با استفاده از

$$\tilde{z}_j = \left(\frac{\hat{m}_{\beta_j^{(d)}} - \hat{m}_{\beta_j^B}}{se(\hat{m}_{\beta_j^{(d)}})}, \frac{\hat{l}_{\beta_j^{(d)}} - \hat{l}_{\beta_j^B}}{se(\hat{l}_{\beta_j^{(d)}})}, \frac{\hat{r}_{\beta_j^{(d)}} - \hat{r}_{\beta_j^B}}{se(\hat{r}_{\beta_j^{(d)}})} \right), j = 0, 1, \dots, p, \quad d = 1, \dots, B$$

استاندارد نمایید.

گام ۳: $\frac{\gamma}{1-\gamma}$ امین و $(1 - \frac{\gamma}{1-\gamma})$ امین چندک را برای مرکز و پهناهای عدد فازی استاندارد

شده $\tilde{z}_k$ محاسبه و به ترتیب با $t_{n-1, 1-\frac{\gamma}{1-\gamma}}^{(*)}$ و $t_{n-1, \frac{\gamma}{1-\gamma}}^{(*)}$ ، $* = m, l, r$ مشخص کنید.

گام چهارم: فاصله اطمینان t - بوت استری متناظر با مرکز و پهناهای چپ و راست را به صورت زیر محاسبه نمایید

$$\begin{cases} \left[\hat{m}_{\beta_j^B} + t_{n-1, \frac{\gamma}{1-\gamma}}^{(m)} se(\hat{m}_{\beta_j^{(d)}}), \hat{m}_{\beta_j^B} + t_{n-1, 1-\frac{\gamma}{1-\gamma}}^{(m)} se(\hat{m}_{\beta_j^{(d)}}) \right] \\ \left[\hat{l}_{\beta_j^B} + t_{n-1, \frac{\gamma}{1-\gamma}}^{(l)} se(\hat{l}_{\beta_j^{(d)}}), \hat{l}_{\beta_j^B} + t_{n-1, 1-\frac{\gamma}{1-\gamma}}^{(l)} se(\hat{l}_{\beta_j^{(d)}}) \right] \\ \left[\hat{r}_{\beta_j^B} + t_{n-1, \frac{\gamma}{1-\gamma}}^{(r)} se(\hat{r}_{\beta_j^{(d)}}), \hat{r}_{\beta_j^B} + t_{n-1, 1-\frac{\gamma}{1-\gamma}}^{(r)} se(\hat{r}_{\beta_j^{(d)}}) \right] \end{cases}$$

در اینجا فرض کرده‌ایم که $\hat{\beta}_j = (\hat{m}_{\beta_j}, \hat{l}_{\beta_j}, \hat{r}_{\beta_j})_{LR}$ برآورد $\hat{\beta}_j = (m_{\beta_j}, l_{\beta_j}, r_{\beta_j})_{LR}$ و $se(\hat{\beta}_j)$ و $\tilde{\beta}_j = (m_{\beta_j}, l_{\beta_j}, r_{\beta_j})_{LR}$ بیانگر انحراف معیار آن باشد. همچنین $t_{n-1, 1-\frac{\gamma}{1-\gamma}}^{(*)}$ و $t_{n-1, \frac{\gamma}{1-\gamma}}^{(*)}$ ، $* = m, l, r$ به ترتیب نقاط بحرانی پایین و بالای سطح γ ، از توزیع t با $n-1$ درجه آزادی هستند.

۴.۴ نمایش‌های عددی

در این بخش، ما یک مطالعه عددی گسترده برای ارزیابی عملکرد مدل پیشنهادی خود انجام می‌دهیم که از دو بخش شامل مثال‌های شبیه‌سازی و واقعی تشکیل شده است. در مطالعات شبیه‌سازی، مقادیر واقعی ضرایب رگرسیونی، شناخته شده است. اما در داده‌های واقعی، علاوه بر انتخاب متغیرها توسط برآوردهای جریمه شده، ما همچنین خطای استاندارد ضرایب برآورد شده در مدل‌های پیشنهادی و مدل‌های دیگر رقیبان را بر اساس نمونه‌های بوت استرپ ارزیابی می‌کنیم (برای اطلاعات بیشتر به [۴۸، ۴۶، ۳۶] مراجعه نمایید).

در حقیقت، از این عامل نایقینی با استفاده از روش بوت استرپ، برای ارزیابی خطای استاندارد برآورد ضرایب رگرسیون استفاده می‌شود. علاوه بر این، در داده‌های واقعی، به

منظور انتخاب متغیر، ما CI های t -بوت استرپی را برای ارزیابی ضرایب برآورد شده ارائه می‌دهیم. در این مسیر، ما نتایج خود را با استفاده از معیارهای بهینگی مدل (GOF) با مدل‌های [۱۵۰، ۸۱، ۳۶] مقایسه می‌کنیم. علاوه بر آن، برای داده‌های واقعی، معیار نایقینی را برای همه مدل‌های مورد بررسی، محاسبه و مورد ارزیابی قرار می‌دهیم. در اینجا مدل‌های ارائه شده توسط کویی و همکاران [۳۶]، کلکین نما و طاهری [۸۱] و مدل زینگ و همکاران [۱۵۰] را به ترتیب با عناوین CDGS، KT و ZFL را نشان می‌دهیم.

تبصره ۱.۴.۴. برای اعداد فازی متقارن، پارامترهای برآورد شده بر اساس معادلات (۱۰.۴) تا (۱۳.۴) به شرح زیر است:

$$\begin{aligned}\hat{\mathbf{h}} &= ((1 + \frac{1}{\varphi} k^2 a^2) \mathbf{X}^T \mathbf{X} + \frac{\lambda}{\varphi} \mathbf{Q}_0) (\mathbf{X}^T \mathbf{y} + \frac{1}{\varphi} k^2 a \mathbf{X}^T \mathbf{s} - \frac{1}{\varphi} k^2 ab \mathbf{X}^T \mathbf{1}) \\ \hat{a} = \hat{c} &= (\mathbf{h}^T \mathbf{X}^T \mathbf{X} \mathbf{h})^{-1} (\mathbf{s}^T \mathbf{X} \mathbf{h} - b \mathbf{1}^T \mathbf{X} \mathbf{h}) \\ \hat{b} = \hat{f} &= \frac{1}{n} (\mathbf{1}^T \mathbf{s} - a \mathbf{1}^T \mathbf{X} \mathbf{h})\end{aligned}$$

که در آن $k_1 = k_2 = k$ و $\mathbf{r} = \mathbf{s}$ است.

۱.۴.۴ مطالعه شبیه‌سازی

در اینجا، ما سه ساختار مختلف را با تعداد ۹، ۱۲ و ۱۰ متغیر ورودی و به ترتیب اندازه نمونه ۱۵، ۲۰ و ۲۰ را در نظر می‌گیریم. بعلاوه، ما برای هر یک از سناریوها تعداد ۳۰ تکرار شبیه‌سازی داده خواهیم داشت.

• ساختار اول: مدل رگرسیونی فازی زیر را در نظر بگیرید:

$$\hat{\mathbf{y}}_i = \oplus_{j=0}^9 (\tilde{\beta}_j \otimes X_{ij}),$$

طوری که در آن داریم:

$$\begin{aligned}j = 1, 2 \text{ برای } & X_{ij} = z_{ij} + \epsilon_{ij}, z_{ij} \sim N(1/5, 1) \quad , \quad \epsilon_{ij} \sim N(0, 0.02) \\ j = 3, \dots, 6 \text{ برای } & [\Sigma_{lk}] = Cov(X_{il}, X_{ik}) = 0.8^{|l-k|}, (X_{i3}, \dots, X_{i6}) \sim N(0, \Sigma) \\ j = 7, 8, 9 \text{ برای } & X_{ij} = z_{ij} + \epsilon_{ij}, z_{ij} \sim N(0, 0.6) \quad , \quad \epsilon_{ij} \sim N(0, 0.04)\end{aligned}$$

$$\begin{aligned}\tilde{\beta}_1 = \tilde{\beta}_2 = \tilde{\beta}_3 &= (1/5, 0.5)_T \quad , \quad \tilde{\beta}_0 = (1, 0.2)_T \\ \tilde{\beta}_7 = \tilde{\beta}_8 = \tilde{\beta}_9 &= (0, 0)_T \quad , \quad \tilde{\beta}_4 = \tilde{\beta}_5 = \tilde{\beta}_6 = (0, 0.4)_T\end{aligned}$$

$$\tilde{\epsilon}_i = (\epsilon_i, s_{\epsilon_i})_T, s_{\epsilon_i} = |\epsilon_i| \quad , \quad \epsilon_i \sim N(0, 2) \quad .3$$

• ساختار دوم: مدل رگرسیونی فازی زیر را در نظر بگیرید:

$$\hat{\mathbf{y}}_i = \oplus_{j=0}^{12} (\tilde{\beta}_j \otimes X_{ij}),$$

طوری که در آن داریم:

$$\begin{aligned} 1. \quad & j = 1, 2, 3 \text{ برای } X_{ij} = z_{ij} + \epsilon_{ij}, z_{ij} \sim N(0, 0.8) \quad , \quad \epsilon_{ij} \sim N(0, 0.3) \\ & j = 4, \dots, 9 \text{ برای } [\Sigma_{lk}] = Cov(X_{il}, X_{ik}) = 0.9^{|l-k|}, (X_{i4}, \dots, X_{i9}) \sim N(0, \Sigma) \\ & j = 10, 11, 12 \text{ برای } X_{ij} = z_{ij} + \epsilon_{ij}, z_{ij} \sim N(0, 1/2) \quad , \quad \epsilon_{ij} \sim N(0, 0.4) \end{aligned}$$

$$\begin{aligned} 2. \quad & \tilde{\beta}_1 = \tilde{\beta}_2 = \tilde{\beta}_3 = \tilde{\beta}_4 = (1/5, 0.8)_T \quad , \quad \tilde{\beta}_0 = (0, 0)_T \\ & \tilde{\beta}_9 = \tilde{\beta}_{10} = \tilde{\beta}_{11} = \tilde{\beta}_{12} = (0, 0)_T \quad , \quad \tilde{\beta}_5 = \tilde{\beta}_6 = \tilde{\beta}_7 = \tilde{\beta}_8 = (0, 0.5)_T \end{aligned}$$

$$3. \quad \tilde{\epsilon}_i = (\epsilon_i, s_{\epsilon_i})_T, s_{\epsilon_i} = |\epsilon_i| \quad , \quad \epsilon_i \sim N(0, 2)$$

• ساختار سوم: مدل رگرسیونی فازی زیر را در نظر بگیرید:

$$\hat{\mathbf{y}}_i = \bigoplus_{j=0}^{10} (\hat{\beta}_j \otimes X_{ij}),$$

طوری که در آن داریم:

$$\begin{aligned} 1. \quad & j = 1, 2 \text{ برای } X_{ij} = z_{ij} + \epsilon_{ij}, z_{ij} \sim N(1, 1) \quad , \quad \epsilon_{ij} \sim N(0, 0.5) \\ & j = 3, \dots, 7 \text{ برای } [\Sigma_{lk}] = Cov(X_{il}, X_{ik}) = 0.5^{|l-k|}, (X_{i3}, \dots, X_{i7}) \sim N(0, \Sigma) \\ & j = 8, 9, 10 \text{ برای } X_{ij} = z_{ij} + \epsilon_{ij}, z_{ij} \sim N(0, 1) \quad , \quad \epsilon_{ij} \sim N(0, 0.2) \end{aligned}$$

$$\begin{aligned} 2. \quad & \tilde{\beta}_1 = \tilde{\beta}_2 = (2, 1)_T \quad , \quad \tilde{\beta}_0 = (0, 0)_T \\ & \tilde{\beta}_8 = \tilde{\beta}_9 = \tilde{\beta}_{10} = (0, 0)_T \quad , \quad \tilde{\beta}_3 = \tilde{\beta}_4 = \tilde{\beta}_5 = \tilde{\beta}_6 = \tilde{\beta}_7 = (0, 0.8)_T \end{aligned}$$

$$3. \quad \tilde{\epsilon}_i = (\epsilon_i, s_{\epsilon_i})_T, s_{\epsilon_i} = |\epsilon_i| \quad , \quad \epsilon_i \sim N(0, 1/5)$$

در مقایسه مدل جریمه شده با مدل‌های رقیب انتخاب شده بر اساس میانگین اندازه خوبی برازش بر روی داده‌های شبیه‌سازی شده، جدول ۱.۴ برتری مدل پیشنهادی را نشان می‌دهد. بر این اساس، مدل‌های جریمه شده در مقایسه با مدل‌های CDGS، KT و ZFL برای $t \in \{0.5, 1, 2\}$ ، مقادیر کمتری را در معیارهای فازی خوبی برازش (GOF) تولید می‌کنند.

جدول ۱.۴: مقایسه مدل‌های جریمه با مدل‌های رقیب بر اساس میانگین معیارهای ارزیابی بهینگی در مجموعه داده‌های شبیه‌سازی.

سناریوی ۳		سناریوی ۲		سناریوی ۱		Models
S	CV	S	CV	S	CV	
۰/۲۵۸۷	۶/۲۱۵۵	۰/۲۸۴۷	۴/۷۸۴۱	۰/۲۱۰۵	۹/۵۱۱۸	CDGS
۰/۲۷۳۱	۵/۱۷۴۸	۰/۳۷۲۲	۲/۸۹۵۲	۰/۲۸۷۳	۴/۲۳۱۶	KT
۰/۲۷۹۲	۴/۸۹۴۳	۰/۳۸۸۲	۲/۳۰۶۸	۰/۳۱۱۶	۴/۱۵۲۷	ZFL
۰/۳۴۵۴	۳/۶۵۴۹	۰/۴۱۳۳	۱/۱۰۶۵	۰/۳۶۲۲	۳/۲۱۷۸	$t = ۰/۵$
۰/۳۴۲۲	۳/۹۹۷۸	۰/۴۱۲۷۰۳۶	۱/۴۵۰۶	۰/۳۴۱۱	۳/۷۲۳۹	جریمه شده $t = ۱$
۰/۳۴۲۷	۳/۹۴۲۷	۰/۴۱۱۸	۱/۳۱۰۷	۰/۳۴۶۲	۳/۴۳۵۱	$t = ۲$

همچنین، کارایی داده‌های شبیه‌سازی شده ۸، ۱۵ و ۱۱ ام به ترتیب در هر کدام از ساختارهای شبیه‌سازی شده در جدول ۲.۴ به عنوان بهترین مدل جریمه‌ای، قابل مشاهده است.

در جدول ۱.۴، در سناریوی اول، مقادیر به‌دست آمده برای پارامتر تنظیم λ در هر کدام از حالات $t \in \{۰/۵, ۱, ۲\}$ به ترتیب ۱۵/۴۰، ۸/۰۲، ۴/۹۸، در سناریوی دوم این مقادیر برابر با ۹/۲۴، ۳/۶۸، ۱/۳۲، و در سناریوی سوم برابر با ۷/۷۳، ۳/۵۴، ۰/۹۶ شده است.

همچنین انتخاب خودکار متغیرهای موثر توسط مدل‌های جریمه‌ای برای $t = ۰/۵$ و $t = ۱$ (LASSO) را می‌توان مشاهده کرد. به بیانی دیگر، به طور خودکار، در سناریوی اول برای مدل جریمه‌ای با $t = ۰/۵$ ، متغیرهای سوم و هشتم و برای مدل جریمه‌ای با $t = ۱$ (LASSO) عدد ثابت مدل و متغیرهای سوم و چهارم حذف شده‌اند. همین طور در سناریوی دوم نیز برای $t = ۰/۵$ متغیرهای پنجم، ششم و نهم و برای $t = ۱$ (LASSO) متغیرهای ششم، هفتم، هشتم، نهم، دهم و یازدهم به طور خودکار حذف شده‌اند. در سناریوی سوم نیز این حذف خودکار انجام گرفته است به طوری که برای مدل جریمه‌ای با $t = ۰/۵$ متغیرهای چهارم و هشتم و برای حالت $t = ۱$ (LASSO)، عدد ثابت مدل و متغیرهای پنجم، ششم و هشتم از مدل حذف شده‌اند. توجه داشته باشید که در تمام سناریوها مدل جریمه‌ای ریج ($t = ۲$)، هیچ متغیری را از مدل حذف نکرده است.

در جدول ۲.۴ مدل‌های برازش شده در هر سناریو را برای هر سه مدل جریمه‌ای به تفکیک ارائه داده ایم.

جدول ۲.۴: نتایج مدل‌های جریبه‌ی ۸، ۱۵، ۱۱، به ترتیب در سناریوهای اول تا سوم.

Coefficients	سناریوی اول		سناریوی دوم		سناریوی سوم	
	$t = 1$	$t = 2$	$t = 1$	$t = 2$	$t = 1$	$t = 2$
$\hat{\beta}_0$	$(-2/5 \times 10^{-9}, 2/628)_T$	$(0/447, 2/673)_T$	$(-4 \times 10^{-9}, 1/561)_T$	$(-0/007, 1/560)_T$	$(2/3 \times 10^{-9}, 0/668)_T$	$(0/011, 0/005)_T$
$\hat{\beta}_1$	$(2/128, 0/000)_T$	$(2/001, 0/000)_T$	$(1/490, 0/000)_T$	$(1/495, 0/000)_T$	$(1/815, 0/081)_T$	$(1/973, 0/523)_T$
$\hat{\beta}_2$	$(0/840, 0/000)_T$	$(0/000, 0/000)_T$	$(1/484, 0/000)_T$	$(1/494, 0/000)_T$	$(1/334, 0/604)_T$	$(1/683, 0/651)_T$
$\hat{\beta}_3$	$(6/0 \times 10^{-8}, 7/9 \times 10^{-4})_T$	$(0/151, 0/000)_T$	$(1/493, 0/000)_T$	$(1/501, 0/000)_T$	$(0/742, 0, 485)_T$	$(0/671, 0/446)_T$
$\hat{\beta}_4$	$(-0/106, 0/125)_T$	$(0/004, 0/000)_T$	$(1/30, 0/000)_T$	$(1/325, 0/000)_T$	$(0/016, 0/415)_T$	$(0/104, 0/458)_T$
$\hat{\beta}_5$	$(0/498, 0/000)_T$	$(0/569, 0/000)_T$	$(0/153, 0/000)_T$	$(0/257, 0/000)_T$	$(0/009, 0/302)_T$	$(0/017, 0/376)_T$
$\hat{\beta}_6$	$(-0/009, 0/842)_T$	$(-0/477, 0/649)_T$	$(1/4 \times 10^{-9}, 0/000)_T$	$(-0/507, 0/000)_T$	$(0/37, 0/512)_T$	$(0/034, 0/402)_T$
$\hat{\beta}_7$	$(0/172, 0/000)_T$	$(0/674, 0/000)_T$	$(0/014, 0/000)_T$	$(-1/648, 0/000)_T$	$(0/013, 0/287)_T$	$(0/021, 0/183)_T$
$\hat{\beta}_8$	$(0/001, 0/000)_T$	$(-0/307, 0/176)_T$	$(1/3 \times 10^{-4}, 0/000)_T$	$(-0/076, 0/005)_T$	$(5/1 \times 10^{-9}, 0/000)_T$	$(0/012, 0/008)_T$
$\hat{\beta}_9$	$(-1/102, 1/308)_T$	$(-0/803, 1/105)_T$	$(-8 \times 10^{-10}, 6/0 \times 10^{-11})_T$	$(-0/197, 0/012)_T$	$(0/004, 0/048)_T$	$(0/015, 0/047)_T$
$\hat{\beta}_{10}$	—	—	$(3/5 \times 10^{-9}, 2 \times 10^{-10})_T$	$(0/014, 0/000)_T$	$(0/011, 0/163)_T$	$(0/022, 0/034)_T$
$\hat{\beta}_{11}$	—	—	$(5/5 \times 10^{-9}, 2 \times 10^{-10})_T$	$(0/015, 0/000)_T$	—	—
$\hat{\beta}_{12}$	—	—	$(-0/017, 0/004)_T$	$(-0/025, 0/002)_T$	—	—
معیار ارزیابی بهینگی	$1/5372$	$1/6710$	$0/765$	$0/788$	$1/825$	$1/977$
S	$0/5048$	$0/4683$	$0/5162$	$0/5092$	$0/4267$	$0/403$
پارامتر تنظیم	$15/40$	$8/02$	$9/24$	$1/22$	$7/73$	$0/96$
مدل‌های جریبه‌ای	$\hat{y}_i = \oplus_{j=1}^8 (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j=1}^8 (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j=1}^{12} (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j=1}^{12} (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j=1}^{12} (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j=1}^{12} (\hat{\beta}_j \otimes X_{ij})$

۲.۴.۴ بخش مثال‌های واقعی

تبصره ۲.۴.۴. در تمام مثال‌ها $\gamma = 0.95$ در نظر گرفته شده است.

تبصره ۳.۴.۴. در تمام مثال‌های واقعی برای بررسی وجود همخطی در بین متغیرهای پیشگو از شاخص ضریب تورم واریانس^۳ (VIF) با رابطه $VIF_j = \frac{1}{1-R_j^2}$ ، $j = 1, 2, \dots, p$ و آماره کاپا^۴ که به آن عدد شرطی^۵ نیز می‌گویند، با رابطه $\kappa = \sqrt{\frac{\lambda_1}{\lambda_p}}$ استفاده می‌شود. که در آن R_j ، λ_1 و λ_p به ترتیب ضریب همبستگی متغیر j ام با متغیر پاسخ، کوچکترین و بزرگترین مقدار ویژه ماتریس $X^T X$ هستند. مقادیر بالای 10 در این شاخص بیانگر وجود همخطی در متغیرهای پیشگو است [۱۸].

داده‌های مزه پنیر

این داده‌ها که در بخش مثال‌های انگیزشی به عنوان سومین مجموعه داده در جدول ۳.۲ آمده است، برای ارزیابی تأثیر سه متغیر بر مزه پنیر در یک آزمایش واقعی تولید شده اند. در داده‌ها، اسید استیک، سولفید هیدروژن و اسید لاکتیک به ترتیب سه متغیر ورودی هستند. متغیر خروجی نیز به عنوان عطر و طعم پنیر در نظر گرفته می‌شود که توسط یک متخصص به شکل یک عدد فازی مثلثی مقارن ارائه می‌شود. داده‌های این مثال در بخش داده‌های انگیزشی به عنوان دومین مجموعه داده آمده است.

در ابتدا برای تشخیص وجود یا عدم وجود همخطی، شاخص ضریب تورم واریانس را محاسبه کردیم و که مقادیر زیر برای متغیرهای سه گانه به ترتیب $VIF_1 = 12/064$ ، $VIF_2 = 7/914$ و $VIF_3 = 2/808$ به دست آمده است. مقدار عددی بالای 10 وجود همخطی را نشان می‌دهد [۵۴]. علاوه بر این، مقدار عددی شاخص کاپا $\sqrt{\lambda_1/\lambda_3} = 13/67$ است که در آن λ_1 و λ_3 به ترتیب کوچکترین و بزرگترین مقادیر ویژه ماتریس $X^T X$ هستند. این مقادیر و معیارها وجود همخطی خطی قوی را که در ماتریس طرح تأیید می‌کنند. نتیجه برازش مدل جریمه شده با $t \in \{0.5, 1, 2\}$ و سه رقیب مورد مطالعه آن در جداول ۳.۴ و ۴.۴ خلاصه شده است. بدیهی است که روش جریمه شده برای $t = 0.5$ و $t = 1$ (LASSO) متغیرهای مهم را انتخاب کرده اند. از این رو، آن‌ها به ترتیب با حذف متغیرهای سوم (اسید لاکتیک) و اول (اسید استیک)، فرایند انتخاب متغیر را به طور خودکار انجام داده‌اند.

^۳Variance inflation factor

^۴Kappa

^۵Condition Number

جدول ۳.۴: معیارهای ارزیابی مدل جهت مقایسه مدل‌های جریمه و مدل‌های رقیب

معیار ارزیابی بهینگی			مدل
CV	S	P	
۱۳/۸۵۹	۰/۱۰۹	۱۴/۴۷۵	CDGS
۷/۶۴۸	۰/۲۴۵	۸/۲۹۶	KT
۷/۰۱۷	۰/۲۸۱	۷/۳۲۴	ZFL
۶/۰۵۴	۰/۳۰۴	۶/۲۳۲	$t = ۰/۵$
۶/۸۸۹	۰/۲۹۳	۷/۲۷۶	جریمه شده $t = ۱$
۷/۵۳۰	۰/۲۵۱	۸/۲۷۶	$t = ۲$

با توجه به نتایج جدول ۳.۴ کاهش قابل توجه مقادیر P و CV در مدل‌های جریمه‌ای (خصوصاً در حالت‌های $t = ۰/۵, ۱$) و افزایش شاخص شباهت مدل در آن‌ها در مقابل مدل‌های رقیب نشان از عملکرد مطلوب‌تر آن‌ها در برآزش مدل است. این درحالی است که ما حذف خودکار متغیر را نیز در آن‌ها داشته ایم.

جدول ۴.۴: ضرایب مدل‌های جریمه‌ای برای داده‌های مزه پنیر

ضرایب				مدل
$\hat{\beta}_3$	$\hat{\beta}_2$	$\hat{\beta}_1$	$\hat{\beta}_0$	
$(۲/۷۶۴, ۰/۴۱۵)_T$	$(-۲/۹۱۹, ۰/۰۰۰)_T$	$(۳۱/۱۱۵, ۴/۶۶۷)_T$	$(-۱۲۷/۶۹۳, ۰/۰۰۰)_T$	CDGS
$(۳۰/۲۰۵, ۰/۴۱۶)_T$	$(۵/۴۹۶, ۰/۰۰۰)_T$	$(-۵/۴۵۹, ۰/۴۲۲)_T$	$(-۱۶/۹۶۷, ۶/۵۷۲)_T$	KT
$(۳۱/۷۲۷, ۱/۱۳۵)_T$	$(۵/۳۸۰, ۰/۹۳۳)_T$	$(-۵/۱۱۳, ۰/۰۰۰)_T$	$(-۲۰/۴۴۹, ۰/۰۰۰)_T$	ZFL
$(-۲/۵ \times ۱۰^{-۹}, ۰/۰۰۰)_T$	$(-۲/۲۳۷, ۰/۰۰۰)_T$	$(۲۹/۹۵۵, ۰/۶۹۱)_T$	$(-۱۲۱/۳۹۱, ۱/۶۴۶)_T$	$t = ۰/۵$
$(۲۰/۳۷۴, ۱/۷۰۶)_T$	$(۳/۷۵۲, ۰/۳۱۴)_T$	$(-۲/۶ \times ۱۰^{-۹}, ۰/۰۰۰)_T$	$(-۲۲/۴۸۲, ۰/۰۰۰)_T$	$(\lambda = ۳۹/۰۱)$
				$t = ۱$
				جریمه شده $(\lambda = ۲۴)$
$(۱۴/۵۵۴, ۲/۴۷۳)_T$	$(۳/۹۰۵, ۰/۶۶۴)_T$	$(۲/۰۸۹, ۰/۳۵۵)_T$	$(-۲۵/۸۱۰, ۰/۰۰۰)_T$	$t = ۲$
				$(\lambda = ۰/۵۰)$

جدول ۴.۴ ضرایب مدل‌های جریمه‌ای و سه رقیب دیگر را ارائه می‌کند و برای مثال مدل رگرسیونی فازی پیشنهاد شده توسط مدل جریمه‌ای LASSO به شرح زیر است:

$$\hat{y}_i = \bigoplus_{j \neq 1}^3 (\hat{\beta}_j \otimes X_{ij}), \quad t = 1, \quad \lambda = ۳۹/۰۱$$

که در آن $\hat{y}$ خروجی برآورد شده مدل فازی است و x_2 و x_3 که به ترتیب سولفید هیدروژن H_2S و اسیدلاکتیک هستند به عنوان متغیرهای ورودی در مدل باقی مانده‌اند.

اکنون ما فواصل اطمینان t -بوت استری و خطاهای استاندارد (se) را برای بررسی صحت ضرایب برآورد شده در مدل‌های فوق‌الذکر، به ترتیب در جداول ۶.۴ و ۵.۴ گزارش کرده‌ایم. توجه داشته باشید که فواصل اطمینان t -بوت استری و se از الگوریتم ۵ بدست آمده است

که براساس روش پیشنهادی [۴۸] است. در تجزیه و تحلیل عددی، ۸۰۰ نمونه بوت استریپی استفاده شده است.

جدول ۵.۴: فاصله اطمینان t - بوت استریپی برای پارامترها در مدل‌های مختلف (اعداد داخل پرانتز طول فاصله اطمینان است).

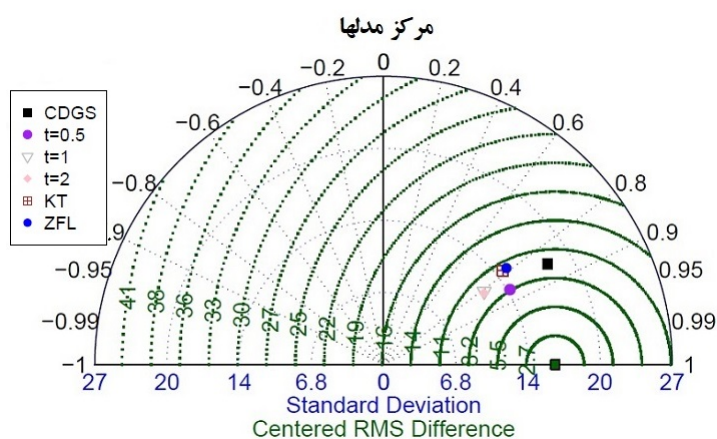
مقدار t فاصله اطمینان	$\tilde{\beta}_0$	$\tilde{\beta}_1$	$\tilde{\beta}_2$	$\tilde{\beta}_3$
مرکز ۰/۵	$[-۱۲۲/۸۴۵, -۱۱۹/۷۴۰]$ (۳/۱۰۵)	$[۲۹/۲۹۳, ۳۱/۹۴۷]$ (۲/۶۵۴)	$[-۳/۷۳۱, -۲/۱۹۰]$ (۱/۵۴۱)	$[-۰/۰۴۹, ۰/۰۲۳]$ (۰/۰۷۲)
	$[۱/۴۹۷, ۱/۷۸۰]$ (۰/۲۸۳)	$[۰/۵۲۲, ۰/۸۷۴]$ (۰/۳۵۲)	$[۰/۰۰۰, ۰/۰۶۸]$ (۰/۰۶۸)	$[۰/۰۰۰, ۰/۰۲۰]$ (۰/۰۲۰)
	$[-۲۳/۸۶۲, -۲۱/۰۲۵]$ (۲/۸۳۷)	$[-۰/۰۲۱, ۰/۰۱۹]$ (۰/۰۴۰)	$[۲/۸۳۴, ۴/۵۰۴]$ (۱/۶۷۰)	$[۱۸/۶۹۶, ۲۱/۳۶۸]$ (۲/۶۷۲)
	$[۰/۰۰۰, ۰/۱۴۲]$ (۰/۱۴۲)	$[۰/۰۰۰, ۰/۰۱۹]$ (۰/۰۱۹)	$[۰/۲۰۹, ۰/۵۲۱]$ (۰/۳۱۲)	$[۱/۴۴۲, ۱/۸۲۶]$ (۰/۳۸۴)
مرکز ۱	$[-۲۳/۳۹۱, -۲۳/۲۶۹]$ (۴/۱۲۲)	$[۰/۵۵۵, ۲/۶۹۷]$ (۲/۱۴۲)	$[۱/۱۶۲, ۳/۶۸۵]$ (۲/۵۲۳)	$[۱۲/۲۱۵, ۱۵/۴۳۳]$ (۳/۲۱۸)
	$[۰/۰۰۰, ۰/۱۳۶]$ (۰/۱۳۶)	$[۰/۰۹۰, ۰/۴۷۳]$ (۰/۳۸۳)	$[۰/۳۵۹, ۰/۷۶۲]$ (۰/۴۰۳)	$[۲/۱۱۳, ۲/۵۴۴]$ (۰/۴۳۱)
	$[-۱۳۵/۰۷۲, -۱۲۲/۲۹۱]$ (۱۲/۷۸۱)	$[۲۵/۰۸۴, ۳۵/۱۷۲]$ (۱۰/۸۸)	$[-۴/۸۶۱, ۱/۵۳۸]$ (۶/۳۹۹)	$[-۰/۴۵۴, ۶/۳۶۸]$ (۶/۸۲۲)
	$[۰/۰۰۰, ۰/۶۴۷]$ (۰/۶۴۷)	$[۲/۹۳۸, ۵/۳۸۹]$ (۲/۴۵۱)	$[۰/۰۰۰, ۰/۷۰۵]$ (۰/۷۰۵)	$[۰/۰۰۰, ۰/۸۷۹]$ (۰/۸۷۹)
مرکز ۲	$[-۲۱/۸۷۱, -۱۳/۴۷۹]$ (۸/۳۹۲)	$[-۸/۱۸۷, -۱/۹۳۴]$ (۶/۲۵۳)	$[۳/۲۵۷, ۷/۷۸۹]$ (۴/۵۳۲)	$[۲۸/۵۹۱, ۳۳/۸۲۳]$ (۵/۲۳۲)
	$[۶/۰۵۷, ۷/۱۷۹]$ (۱/۱۲۲)	$[۰/۰۰۰, ۰/۹۵۵]$ (۰/۹۵۵)	$[۰/۰۰۰, ۰/۵۷۴]$ (۰/۵۷۴)	$[۰/۰۰۰, ۰/۷۱۹]$ (۰/۷۱۹)
	$[-۲۴/۷۲۹, -۱۷/۹۸۵]$ (۶/۷۴۴)	$[-۷/۷۸۹, -۳/۱۹۶]$ (۴/۵۹۳)	$[۳/۳۹۵, ۷/۴۸۲]$ (۴/۰۸۷)	$[۲۹/۹۵۸, ۳۴/۵۷۰]$ (۴/۶۱۲)
	$[۰/۰۰۰, ۰/۴۸۴]$ (۰/۴۸۴)	$[۰/۰۰۰, ۰/۶۵۲]$ (۰/۶۵۲)	$[۰/۴۳۵, ۱/۲۵۹]$ (۰/۸۲۴)	$[۰/۵۷۱, ۱/۵۴۳]$ (۰/۹۷۲)
مرکز -	$[-۱۳۵/۰۷۲, -۱۲۲/۲۹۱]$ (۱۲/۷۸۱)	$[۲۵/۰۸۴, ۳۵/۱۷۲]$ (۱۰/۸۸)	$[-۴/۸۶۱, ۱/۵۳۸]$ (۶/۳۹۹)	$[-۰/۴۵۴, ۶/۳۶۸]$ (۶/۸۲۲)
	$[۰/۰۰۰, ۰/۶۴۷]$ (۰/۶۴۷)	$[۲/۹۳۸, ۵/۳۸۹]$ (۲/۴۵۱)	$[۰/۰۰۰, ۰/۷۰۵]$ (۰/۷۰۵)	$[۰/۰۰۰, ۰/۸۷۹]$ (۰/۸۷۹)
	$[-۲۱/۸۷۱, -۱۳/۴۷۹]$ (۸/۳۹۲)	$[-۸/۱۸۷, -۱/۹۳۴]$ (۶/۲۵۳)	$[۳/۲۵۷, ۷/۷۸۹]$ (۴/۵۳۲)	$[۲۸/۵۹۱, ۳۳/۸۲۳]$ (۵/۲۳۲)
	$[۶/۰۵۷, ۷/۱۷۹]$ (۱/۱۲۲)	$[۰/۰۰۰, ۰/۹۵۵]$ (۰/۹۵۵)	$[۰/۰۰۰, ۰/۵۷۴]$ (۰/۵۷۴)	$[۰/۰۰۰, ۰/۷۱۹]$ (۰/۷۱۹)
مرکز -	$[-۲۴/۷۲۹, -۱۷/۹۸۵]$ (۶/۷۴۴)	$[-۷/۷۸۹, -۳/۱۹۶]$ (۴/۵۹۳)	$[۳/۳۹۵, ۷/۴۸۲]$ (۴/۰۸۷)	$[۲۹/۹۵۸, ۳۴/۵۷۰]$ (۴/۶۱۲)
	$[۰/۰۰۰, ۰/۴۸۴]$ (۰/۴۸۴)	$[۰/۰۰۰, ۰/۶۵۲]$ (۰/۶۵۲)	$[۰/۴۳۵, ۱/۲۵۹]$ (۰/۸۲۴)	$[۰/۵۷۱, ۱/۵۴۳]$ (۰/۹۷۲)
	$[-۱۳۵/۰۷۲, -۱۲۲/۲۹۱]$ (۱۲/۷۸۱)	$[۲۵/۰۸۴, ۳۵/۱۷۲]$ (۱۰/۸۸)	$[-۴/۸۶۱, ۱/۵۳۸]$ (۶/۳۹۹)	$[-۰/۴۵۴, ۶/۳۶۸]$ (۶/۸۲۲)
	$[۰/۰۰۰, ۰/۶۴۷]$ (۰/۶۴۷)	$[۲/۹۳۸, ۵/۳۸۹]$ (۲/۴۵۱)	$[۰/۰۰۰, ۰/۷۰۵]$ (۰/۷۰۵)	$[۰/۰۰۰, ۰/۸۷۹]$ (۰/۸۷۹)

بر اساس جدول ۵.۴، در مدل‌های جرمیه شده طول فاصله اطمینان‌های به‌دست آمده یعنی مقادیر داخل پرانتز، برای ضرایب مدل هم در مراکز و هم در پهناها در مقایسه با سایر مدل‌ها بسیار کوتاه‌تر است. علاوه بر آن، می‌دانیم که اگر حد پایین فاصله اطمینانی منفی باشد، به اجبار آن را صفر در نظر می‌گیریم که این وضعیت برای پهناها در مدل‌های رقیب به دفعات رخ داده است. این کار باعث کوتاه‌تر شدن طول این فاصله‌ها می‌شود، با این حال مدل‌های جرمیه شده در مقایسه با سایر مدل‌ها بدلیل کوتاه‌تر بودن فواصل اطمینانشان برتر هستند.

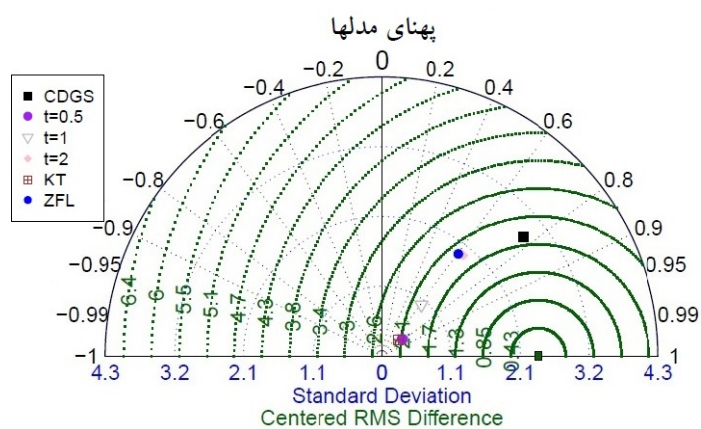
همچنین نتایج بدست آمده در جدول ۶.۴ از طریق محاسبه مقادیر se برای تمامی پارامترهای مدل‌های مورد آزمایش، نشان می‌دهد که مدل‌های جرمیه‌ای در مقایسه با مدل‌های رقیب دارای se کمتری برای ضرایب خود هستند. لذا، میزان نایقینی از صحت درستی برآورد ضرایب در این مدل‌ها کمتر است که این امر به‌واسطه خاصیت ذاتی موجود در مدل‌های جرمیه‌ای است.

جدول ۶.۴: مقدار se برای ضرایب مدل‌ها در داده‌های مزه پنی

پهنا				مرکز				مدل
$se_{\hat{\beta}_T}$	$se_{\hat{\beta}_Y}$	$se_{\hat{\beta}_1}$	$se_{\hat{\beta}_0}$	$se_{\hat{\beta}_T}$	$se_{\hat{\beta}_Y}$	$se_{\hat{\beta}_1}$	$se_{\hat{\beta}_0}$	
۰/۰۹۶	۰/۲۰۴	۰/۱۷۵	۰/۱۴۳	۱/۱۶۴	۰/۷۳۵	۲/۱۲۲	۲/۷۴۴	CDGS
۰/۰۷۷	۰/۰۹۸	۰/۱۴۶	۰/۱۳۳	۰/۶۰۵	۰/۷۱۶	۰/۹۵۹	۱/۵۳۲	KT
۰/۰۸۳	۰/۰۸۶	۰/۱۳۴	۰/۱۲۹	۰/۷۰۲	۰/۸۸۰	۰/۷۹۴	۱/۲۶۰	ZFL
۰/۰۱۳	۰/۰۴۸	۰/۰۹۶	۰/۰۷۵	۰/۰۲۸	۰/۳۷۷	۰/۴۲۳	۰/۷۹۲	$t = ۰/۵$
								$(\lambda = ۳۹/۰۱)$
۰/۰۶۶	۰/۰۵۵	۰/۰۱۰	۰/۰۸۲	۰/۵۱۲	۰/۴۰۲	۰/۰۱۲	۰/۸۰۳	$t = ۱$
								$(\lambda = ۲۴)$
۰/۰۳۸	۰/۰۵۱	۰/۰۶۵	۰/۰۷۸	۰/۴۱۸	۰/۳۶۴	۰/۳۲۵	۰/۸۱۰	$t = ۲$
								$(\lambda = ۰/۵۰)$



(ا)



(ب)

شکل ۱.۴: مقایسه شهودی-استنباطی مدل‌های جریمه‌ای با مدل‌های رقیب. (ا) مقدار مرکز (ب) مقدار پهنا.

در نمودار تیلور ۱.۴، بررسی مقایسه پاسخ‌های برآورد شده توسط مدل‌های جریمه‌شده و مدل‌های رقیب از منظر همبستگی و تطابق با پاسخهای مشاهده شده، انجام شده است. با توجه به این نمودار، مطابق قسمت (الف)، روش‌های جریمه‌ای پیشنهاد شده برای تمامی مقادیر t از انحراف معیار و میانگین خطای استاندارد کمتری در مقابل مدل‌های رقیب برخوردار هستند. علاوه بر آن، ضریب همبستگی پاسخ‌های برآورد شده با مشاهده شده در مدل‌های جریمه‌ای به‌ویژه برای $t = 0.5$ ، بیشتر از مدل‌های رقیب است. همچنین بر اساس نتایج بخش (ب)، مدل جریمه‌ای لاسو، $t = 0.5$ ، نسبت به دیگر مدل‌ها عملکرد بهتری دارد.

داده‌های قیمت خانه در شانگهای (حالت اعداد فازی نامتقارن)

داده‌های بکار رفته در این مثال با عنوان مجموعه داده‌های چهارم در بخش مثال‌های انگیزشی در جدول ۴.۲ آمده است. در این مثال، برتری مدل‌های مجازات شده پیشنهادی در یک مجموعه داده بزرگ نشان داده می‌شود. این داده‌ها ابتدا توسط [۱۵۱] معرفی شدند. آن‌ها قیمت خرید قابل قبول خانه را به عنوان یک عدد فازی مثلثی نامتقارن با شش متغیر توضیحی غیرفازی در نظر گرفتند که در آن، x_1 ، اندازه مسکن، x_2 ، نرخ بهره وام مسکن، x_3 ، مالیات بر املاک و مستغلات، x_4 ، نسبت پیش پرداخت، x_5 ، درآمد سالانه خانوار و در نهایت x_6 ، جمعیت خانواده را به عنوان متغیرهای ورودی در نظر گرفتند. بر خلاف [۱۵۱]، ما هیچ مشاهده‌ای را حذف نمی‌کنیم و همانطور که در جدول ۴.۲ در بخش مثال‌های انگیزشی به عنوان مجموعه داده چهارم ارائه شده است با کل داده‌ها کار می‌کنیم. توجه داشته باشید که در فرمول‌های برآورد پارامترها برای اعداد فازی مثلثی داریم $k_1 = k_2 = \frac{1}{2}$ باز هم در اینجا ابتدا به بررسی وجود یا عدم وجود همخطی در بین متغیرهای ورودی از طریق محاسبه VIF ها و شاخص $\kappa = \sqrt{\frac{\lambda_1}{\lambda_p}}$ می‌پردازیم. از این رو مقادیر زیر برای VIF های متناظر با هر متغیر به دست آمده است. $VIF_1 = 38/410$ ، $VIF_2 = 31/455$ ، $VIF_3 = 1/970$ ، $VIF_4 = 17/037$ ، $VIF_5 = 17/322$ و $VIF_6 = 18/079$ همچنین شاخص نسبتی زیر نیز که به عنوان معیاری برای وجود همخطی در متغیرهای ورودی است در این مثال برابر $\sqrt{\lambda_1/\lambda_6} = 17/07$ است که باینگر وجود همخطی بین متغیرهای پیشگو است و در آن λ_1 و λ_6 به ترتیب کمترین و بیشترین مقادیر ویژه ماتریس $X^T X$ هستند.

در جدول ۷.۴، معیارهای عملکرد بهینگی مدل برای مدل‌های جریمه شده و مدل‌های رقیب خلاصه شده است. در اینجا علاوه بر MPE، به ترتیب شاخص‌های E_M و E_S را به ترتیب برای اندازه‌گیری اختلاف میانگین بین مراکز پاسخ‌های مشاهده شده و تخمین زده شده و اختلاف بین پهنای پاسخ‌های مشاهده شده و تخمین زده شده در نظر گرفته‌ایم.

مطابق با نتایج جدول ۷.۴، عملکرد مدل‌های جریمه شده در مقابل مدل‌های رقیب بهتر بوده است. این امر ناشی از کمتر بودن مقدار معیار بهینگی برازش مدل که در این جدول ۷.۴ ارائه شده می‌باشد. همچنین، برای مقادیر $t = 0.5$ و $t = 1$ ، مدل‌های مجاز شده پیشنهادی، متغیر چهارم را به عنوان متغیر غیر موثر حذف کرده‌اند این در حالی است که مدل‌های رقیب شامل

جدول ۷.۴: نتایج ارزیابی مدل‌های جریمه شده و مدل‌های رقیب به همراه ضرایب مدل.

مدل‌های دیگر			مدل جریمه شده			ضرایب
ZFL	KT	CDGS	$t = 2$	$t = 1$	$t = 0.5$	
$(-50.61, 0.00, 0.00)_T$	$(-31.70, 0.00, 0.00)_T$	$(-52.71, 0.00, 0.00)_T$	$(-49.33, 0.00, 0.00)_T$	$(-52.25, 0.00, 0.00)_T$	$(-51.46, 0.00, 0.00)_T$	$\hat{\beta}_0$
$(7.179, 0.632, 0.360)_T$	$(7.015, 0.632, 0.395)_T$	$(6.794, 1.933, 1.074)_T$	$(6.644, 0.00, 0.558)_T$	$(6.740, 0.00, 0.152)_T$	$(6.688, 0.00, 0.236)_T$	$\hat{\beta}_1$
$(-58.79, 0.00, 0.00)_T$	$(-56.08, 0.00, 0.00)_T$	$(-4.800, 0.00, 0.00)_T$	$(-7.648, 1.983, 0.000)_T$	$(-5.683, 1.559, 0.000)_T$	$(-5.547, 1.384, 0.000)_T$	$\hat{\beta}_2$
$(-55.52, 0.00, 0.00)_T$	$(-40.67, 0.123, 0.000)_T$	$(-58.29, 0.00, 0.00)_T$	$(-60.20, 1.561, 0.000)_T$	$(-57.93, 1.589, 0.000)_T$	$(-59.00, 1.475, 0.000)_T$	$\hat{\beta}_3$
$(-5.125, 0.00, 0.00)_T$	$(-6.161, 0.00, 0.00)_T$	$(0.096, 0.027, 0.015)_T$	$(-0.193, 0.05, 0.000)_T$	$(-2 \times 10^{-10}, 0.00, 0.00)_T$	$(-4 \times 10^{-9}, 0.00, 0.00)_T$	$\hat{\beta}_4$
$(14.669, 3.158, 2.035)_T$	$(14.503, 3.156, 1.924)_T$	$(14.021, 3.989, 2.217)_T$	$(120.8, 0.00, 0.227)_T$	$(13.982, 0.00, 0.277)_T$	$(13.91, 0.00, 2.571)_T$	$\hat{\beta}_5$
$(-16.48, 0.000, 0.00)_T$	$(-15.43, 0.00, 0.000)_T$	$(2.011, 0.592, 0.329)_T$	$(120.8, 0.00, 0.277)_T$	$(1.622, 0.00, 0.277)_T$	$(1.293, 0.00, 0.239)_T$	$\hat{\beta}_6$
0.993	0.344	66.673	44.197	41.777	41.146	معیارهای ارزیابی بهینگی
$-$	$-$	$-$	0.30	73.80	13.98	پارامتر تنظیم کننده
0.1795	0.135	36.368	36.507	35.726	36.193	میانگین اختلاف مراکز برآورد شده
0.417	0.416	10.233	8.599	8.165	8.480	میانگین اختلاف بهن‌های برآورد شده
$\hat{y}_i = \oplus_{j=0}^6 (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j=0}^6 (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j=0}^6 (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j=0}^6 (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j \neq 4}^6 (\hat{\beta}_j \otimes X_{ij})$	$\hat{y}_i = \oplus_{j \neq 4}^6 (\hat{\beta}_j \otimes X_{ij})$	مدل‌های جریمه شده

جدول ۹.۴: فاصله اطمینان t - بوت استری برای پارامترها در مدل‌های مختلف (اعداد داخل پرانتز طول پرانتز طول فاصله اطمینان است).

روش	مقدار t	فاصله اطمینان	$\hat{\beta}_0$	$\hat{\beta}_1$	$\hat{\beta}_2$	$\hat{\beta}_3$	$\hat{\beta}_4$	$\hat{\beta}_5$	$\hat{\beta}_6$
جریمه شده	۵	مرکز	[۰/۸۳, ۱/۵۸] (۰/۷۵)	[۱۳/۰۷, ۱۴/۵۶] (۱/۴۹)	[-۰/۳۲, ۰/۰۱۵] (۰/۰۵)	[-۶/۲۲, -۵/۶۶] (۴/۵۶)	[-۶/۷۴, -۴/۶۸] (۲/۰۶)	[۴/۲۷, ۸/۷۸] (۴/۵۱)	[-۵/۸۵, -۵/۰۹۵] (۷/۵۶)
		پهنای چپ	[۰/۰۰۰, ۰/۵۲] (۰/۵۲)	[۰/۰۰۰, ۰/۶۱] (۰/۶۱)	[۰/۰۰, ۰/۰۹] (۰/۰۹)	[۱۴/۱۳, ۱۵/۲۴] (۱/۱۱)	[۰/۹۷, ۲/۰۶] (۱/۰۹)	[۰/۰۰, ۰/۸۵] (۰/۸۵)	[۱۴/۸۶, ۱۵/۲۱] (۳/۴۷)
		پهنای راست	[۰/۰۲, ۰/۵۷] (۰/۵۵)	[۱۴/۲, ۲/۸۶] (۱/۴۴)	[۰/۰۰, ۰/۱۱] (۰/۱۱)	[۰/۰۰, ۰/۱۸] (۰/۱۸)	[۰/۰۰, ۰/۰۶۶] (۰/۰۶۶)	[۰/۸۹, ۱/۴۶] (۰/۸۳)	[۰/۰۰, ۰/۳۲] (۰/۳۲)
	۱	مرکز	[۱/۱۷, ۲/۱۳] (۰/۹۶)	[۱۳/۲۲, ۱۴/۵۴] (۱/۳۳)	[۰/۰۰, ۰/۰۴] (۰/۰۴)	[-۵/۹۸, -۵/۴۵] (۵/۳۴)	[-۶/۶۵, -۴/۳۳] (۲/۳۳)	[۴/۳۳, ۸/۷۱] (۴/۳۸)	[-۵/۲۶, -۵/۱۸۹۳] (۷/۹۸)
		پهنای چپ	[۰/۰۰, ۰/۵۶] (۰/۵۶)	[۰/۰۰, ۰/۷۳] (۰/۷۳)	[۰/۰۰, ۰/۰۹] (۰/۰۹)	[۱۵/۱۴, ۱۶/۳۶] (۱/۲۲)	[۱/۳۲, ۲/۳۸] (۱/۰۸)	[۰/۰۰, ۰/۸۲] (۰/۸۲)	[۱۵/۶۴, ۱۶/۰۸۳] (۴/۳۵)
		پهنای راست	[۰/۰۳, ۰/۶۴] (۰/۶۱)	[۰/۰۱, ۰/۵۹] (۰/۵۸)	[۰/۰۰, ۰/۱۳] (۰/۱۳)	[۰/۰۰, ۰/۱۵] (۰/۱۵)	[۰/۰۰, ۰/۱۱] (۰/۱۱)	[۱/۲۵, ۱/۷۷] (۰/۴۲)	[۰/۰۰, ۰/۲۵] (۰/۲۵)
CDGS	۲	مرکز	[۰/۵۲, ۱/۵۶] (۱/۰۴)	[۰/۱۴, ۱/۵۶] (۱/۴۲)	[-۰/۳۳, -۰/۰۹] (۰/۳۴)	[-۶/۲۵, -۵/۱۴] (۵/۴۲)	[-۸/۷۳, -۶/۱۱] (۲/۶۲)	[۴/۵۵, ۹/۷۱] (۵/۱۶)	[-۴/۹۸, -۴/۸۹/۷۹] (۸/۴۴)
		پهنای چپ	[۰/۰۰, ۰/۶۳] (۰/۶۳)	[۰/۰۰, ۰/۸۷] (۰/۸۷)	[۰/۰۰, ۰/۱۶] (۰/۱۶)	[۱۴/۸۹, ۱۶/۱۷] (۱/۲۸)	[۱/۵۳, ۲/۶۴] (۱/۱۱)	[۰/۰۰, ۰/۹۱] (۰/۹۱)	[۱۴/۵۲, ۱۴/۹۶] (۴/۴۲)
		پهنای راست	[۰/۰۳, ۰/۷۱] (۰/۶۸)	[۰/۰۲, ۰/۶۵] (۰/۶۳)	[۰/۰۰, ۰/۱۲] (۰/۱۲)	[۰/۰۰, ۰/۱۷] (۰/۱۷)	[۰/۰۰, ۰/۱۴] (۰/۱۴)	[۱/۰۸, ۱/۵۹] (۰/۵۱)	[۰/۰۰, ۰/۰/۳۸] (۰/۳۸)
	-	مرکز	[-۳/۲۴, ۴/۳۱] (۷/۴۵)	[۱۱/۵۵, ۲/۰/۰] (۹/۱۵)	[-۳/۸۵, ۳/۱۱] (۶/۸۶)	[-۶/۴۳, -۵/۳۴۹] (۱۰/۹۴)	[-۸/۸۴, ۱۲/۰۵] (۹/۰۵)	[-۴/۴۷, ۱۲/۰۵] (۱۶/۵۲)	[-۵/۴۳, -۵/۲۰/۹] (۳/۱۲۶)
		پهنای چپ	[۰/۰۰, ۲/۶۴] (۲/۶۴)	[۱۴/۸, ۵/۸۴] (۴/۱۶)	[۰/۰۰, ۲/۱۳] (۲/۱۳)	[۰/۰۰, ۲/۰۲] (۲/۰۲)	[۰/۰۰, ۱/۵۵] (۱/۵۵)	[۰/۰۰, ۴/۳۸] (۴/۳۸)	[۰/۰۰, ۱/۸۷] (۱/۸۷)
		پهنای راست	[۰/۰۰, ۰/۹۳] (۰/۹۳)	[۰/۵۳, ۴/۱۷] (۳/۶۴)	[۰/۰۰, ۲/۵۷] (۲/۵۷)	[۰/۰۰, ۰/۸۴] (۰/۸۴)	[۰/۰۰, ۱/۵۴] (۱/۵۴)	[۰/۰۰, ۲/۱۳] (۳/۱۳)	[۰/۰۰, ۱/۱۹] (۱/۱۹)
KT	-	مرکز	[-۱۸/۳۰, -۱۲/۳۱] (۵/۸۹)	[۱۰/۳۰, ۱۷/۸۲] (۷/۵۲)	[-۸/۵۷, -۲/۷۳] (۵/۸۴)	[-۴/۶۳, -۳/۵۸۲] (۸/۸۱)	[-۵/۹۴۶, -۵/۷۷۲] (۷/۷۴)	[۱/۸, ۱/۲۱] (۱/۰/۳)	[-۳/۸۸۳, -۲/۴۴۷] (۱۴/۳۶)
		پهنای چپ	[۰/۰۰, ۱/۳۵] (۱/۳۵)	[۱/۸۶, ۵/۱۶] (۳/۳۰)	[۰/۰۰, ۱/۶۵] (۱/۶۵)	[۰/۰۰, ۱/۷۸] (۱/۷۸)	[۰/۰۰, ۱/۳۱] (۱/۳۱)	[۰/۰۰, ۲/۰۶] (۲/۰۶)	[۰/۰۰, ۲/۲۱] (۲/۲۱)
		پهنای راست	[۰/۰۰, ۰/۷۷] (۰/۷۷)	[۰/۵۶, ۲/۴۱] (۱/۸۵)	[۰/۰۰, ۰/۶۷] (۰/۰/۶۷)	[۰/۰۰, ۰/۵۶] (۰/۵۶)	[۰/۰۰, ۱/۰۸] (۱/۰۸)	[۰/۰۰, ۱/۰۳] (۱/۰۳)	[۰/۰۰, ۰/۸۵] (۰/۸۵)
	-	مرکز	[-۱۷/۸, -۱۴/۳۳] (۳/۴۵)	[۱۲/۸۳, ۱۶/۶۴] (۳/۸۱)	[-۶/۶۱, -۲/۸۹] (۳/۲۷)	[-۵/۸۶, -۵/۲۴۳] (۶/۲۱)	[-۶/۳۲, -۵/۵/۰۳] (۸/۱۹)	[۳/۸, ۱/۰/۸۵] (۷/۶۷)	[-۵/۷/۵, -۴/۵۶۴] (۱۱/۵۱)
		پهنای چپ	[۰/۰۰, ۰/۸۲] (۰/۸۲)	[۱/۹۸, ۴/۸۶] (۲/۸۸)	[۰/۰۰, ۱/۲۶] (۱/۲۶)	[۰/۰۰, ۱/۳۳] (۱/۳۳)	[۰/۰۰, ۱/۱۸] (۱/۱۸)	[۰/۰۰, ۰/۹۷] (۰/۹۷)	[۰/۰۰, ۰/۸۴] (۰/۸۴)
		پهنای راست	[۰/۰۰, ۰/۷۲] (۰/۷۲)	[۱/۲۴, ۲/۷۷] (۱/۵۳)	[۰/۰۰, ۰/۶۵] (۰/۶۵)	[۰/۰۰, ۰/۶۲] (۰/۶۲)	[۰/۰۰, ۰/۵۷] (۰/۵۷)	[۰/۰۲, ۰/۷۷] (۰/۸۵)	[۰/۰۰, ۰/۸۲] (۰/۸۲)
ZFL	-	مرکز							
		پهنای چپ							

CDGS، KT و ZFL این ویژگی را ندارند. علاوه بر این، مقادیر ثبت شده برای E_M نشان می‌دهد که میانگین اختلاف بین پاسخ تخمینی و مشاهده شده در مدل‌های جریمه‌ای بطور قابل توجهی کمتر از مدل‌های CDGS، KT و ZFL است. عملکرد بهتر مدل‌های KT و ZFL با توجه به معیار E_S به وضوح عملکرد ضعیف آن‌ها را با توجه در معیار E_M جبران نمی‌کند. بنابراین، عملکرد مدل‌های مجازات پیشنهادی از مدل‌های رقیب بهتر است. اکنون در بررسی فواصل اطمینان و شاخص نایقینی، بر اساس ۱۰۰ نمونه بوت استرپ تولید شده و با استفاده از الگوریتم ۵ و رابطه پیشنهادی در [۴۸] برای فاصله اطمینان‌های بوت استرپ برای ضرایب و se آن‌ها را محاسبه می‌کنیم. همچنین، مقادیر λ بر اساس مدل جریمه شده و ویژگی آن‌ها به طور خودکار بدست آمده است. مقادیر se ضرایب و فاصله اطمینان آن‌ها به ترتیب در جداول ۸.۴ و ۹.۴ ارائه شده است. عملکرد برتر مدل‌های مجازات شده در بررسی نتایج این جداول کاملاً مشهود است.

با توجه به ویژگی انتخاب متغیر در مدل‌های جریمه‌ای این امر باعث کاهش se می‌شوند. این مزیت در فواصل اطمینان‌های ارائه شده نیز مشاهده می‌شود زیرا طول آن‌ها در مقایسه با مدل‌های دیگر کوتاه‌تر است.

۵.۴ نتیجه‌گیری

در تحلیل رگرسیون فازی، انتخاب بهینه متغیرهای ورودی در حضور داده‌های با ابعاد بالا و یا وجود همخطی بین متغیرهای ورودی یک مسئله نظری است که بر عملکرد تحلیل رگرسیون تأثیر می‌گذارد.

با شروع از رویکرد رگرسیون فازی که در مقالات [۴۶، ۴۵، ۴۴، ۴۳، ۴۲، ۳۶] پیشنهاد شده است، ما یک مدل رگرسیون جریمه‌ای فازی با متغیرهای خروجی فازی و متغیرهای ورودی غیر فازی، همراه با راه برآوردهای بهینه پارامتر به صورت فرم‌های بسته با در نظر گرفتن یک روش تخمین کمترین توان دوم جریمه شده پیشنهاد کردیم. مدل رگرسیون فازی پیشنهادی می‌تواند اثرات وجود همخطی در بین متغیرهای ورودی را کنترل کند. علاوه بر این، یک معیار انتخاب متغیر برای انتخاب متغیرهای ورودی (حذف متغیرهای غیر موثر) در روند تهیه مدل رگرسیونی پیشنهاد شده است. از این رو، به منظور بررسی اثربخشی و عملکرد روش پیشنهادی، یک مطالعه شبیه‌سازی تحت سه سناریوی مختلف و دو مثال کاربردی تحت داده‌های واقعی تهیه شد که نتایج حاصله نشان از برتری مدل جریمه‌ای رگرسیون فازی پیشنهادی در مقایسه با سایر مدل‌ها، با توجه به معیارهای نیکویی برازش مدل است. همچنین، برای داده‌های واقعی، برآوردهای به‌دست آمده توسط مدل جریمه شده، دقت بیشتری را در مقایسه با سایر مدل‌ها، با توجه به کم بودن خطاهای استاندارد بوت استرپ و طول بازه‌های اطمینان از خود نشان می‌دهند.

۱.۵.۴ چشم‌اندازهای احتمالی

چشم‌اندازهای احتمالی در این چارچوب نظری عبارتند از:

الف) ترکیب Enet [۱۵۳] و Mnet [۷۳]، رویکردهایی برای انتخاب متغیر که شامل دو جمله جریمه هستند.

ب) گسترش نتایج برای مدل‌سازی در رگرسیون فازی استوار.

ج) بررسی مدل‌های رگرسیون فازی جریمه شده با استفاده از انواع دیگر مترها و فاصله‌ها.

پیوست آ

تعاریف و مفاهیم اولیه

آ.۱ مقدمه

در ریاضیات مدرن، نظریه مجموعه‌های معمولی به‌عنوان پایه و شالوده اصلی آن است. در این نظریه، به گردآیه‌ای معین از اشیاء که دارای یک ویژگی و صفت مطلق هستند مجموعه می‌گویند. ویژگی مورد بحث (فقط در مجموعه‌ها) باید خوش‌تعریف باشد. به‌طور مثال اگر مجموعه مرجع X ، مجموعه اعداد حقیقی باشد و خاصیت P ویژگی «اول بودن»، آن‌گاه P یک ویژگی خوش‌تعریف است که اگر A مجموعه شامل اعضاء این صفت باشد، هر عدد مانند a در اعداد حقیقی یا عضو A است، یعنی اول است و یا عضو A نیست یعنی اول نیست.

اما همواره ویژگی‌هایی که با آن روبرو هستیم، مطلق نیستند، مانند کوتاه بودن، کافی بودن، تقریباً کوچک یا بزرگ بودن و علاوه‌بر آن در علوم رفتاری، روان‌شناسی و حوزه علوم انسانی ما با مفاهیم و تعاریف نادقیقی سروکار داریم که در قالب نظریه مجموعه‌های معمولی جایگاهی ندارند. بدین شکل قالب جدیدی برای پوشش این مفاهیم نادقیق احساس می‌گردد. نظریه مجموعه‌های فازی این امر را عهده‌دار شده و پوشش مناسبی به این مجموعه‌ها ارائه می‌دهد. در واقع تعمیمی طبیعی از نظریه مجموعه‌های معمولی، نظریه مجموعه‌های فازی است. معرفی و تشریح این نظریه در سال ۱۹۶۵ میلادی (۱۳۴۴ ه.ش.) توسط پرفسور لطفی عسگرزاده [۱۴۹] دانشمند ایرانی تبار عرضه شد. این نظریه از زمان ارائه تاکنون،

گسترش و تعمیم زیادی یافته و کاربردهای گوناگونی در زمینه‌های مختلف پیدا کرده است. به‌طور خلاصه، نظریه‌ی مجموعه‌های فازی نظریه‌ای در شرایط عدم اطمینان است. این نظریه قادر است بسیاری از مفاهیم و متغیرها و سیستم‌هایی را که نادقیق و مبهم هستند، همان‌گونه که در عالم واقع اکثراً چنین است، صورت‌بندی ریاضی بخشیده و زمینه را برای استدلال، استنتاج، کنترل و تصمیم‌گیری در شرایط عدم اطمینان فراهم آورد. در ادامه برخی تعاریف اولیه مورد استفاده در مجموعه‌های فازی را می‌آوریم. مطالب این فصل عمدتاً برگرفته از [۲] است.

۲.۱ اصل توسیع و اعداد فازی

در مسیر استفاده از مجموعه‌های فازی و توابع عضویت آن‌ها، امکان بکارگیری و تغییر تابعی از این مجموعه‌ها وجود دارد. در واقع استفاده از عملگرهای مختلف ریاضی اجتناب‌ناپذیر است. از این‌رو نیاز به یک رویکرد برای نحوه عملکرد و ایجاد این تغییرات را خواهیم داشت که در قالب اصل توسیع ارائه می‌گردد.

تعریف ۱.۲. فرض کنید تابع f از X به Y با ضابطه $y = f(x)$ باشد. تابع f به هر نقطه از X ، نقطه‌ای از Y را نسبت می‌دهد. این روش عملیاتی هر تابع مانند f در مجموعه اعداد حقیقی است. حال اگر بخواهیم یک مجموعه (مجموعه‌ی فازی $\tilde{A}$) را تحت تابع f به یک مجموعه (مجموعه‌ی فازی $\tilde{B}$) نسبت دهیم، اصل توسیع را خواهیم داشت.

تعریف ۲.۲ (اصل توسیع برای حالت یک متغیره). فرض کنید X و Y دو مجموعه مرجع و f یک تابع به صورت $f: X \rightarrow Y$ و $\tilde{A}$ یک زیرمجموعه فازی از X باشد. در این صورت $\tilde{B} = f(\tilde{A})$ به صورت یک مجموعه فازی از Y با تابع عضویت زیر تعریف می‌شود،

$$\tilde{B}(y) = \sup_{x, y=f(x)} \tilde{A}(x) \quad (1.A)$$

$\tilde{B}(y) = 0$ است اگر تصویر معکوس y تحت تابع f ، تهی باشد. به تابع f تابع تصویر یا انتقال گوئیم.

تعریف ۳.۲ (اصل توسیع برای حالت چند متغیره). فرض کنید $X_1, X_2, \dots, X_n$ ، n مجموعه مرجع و $X = X_1 \times X_2 \times \dots \times X_n$ حاصل ضرب دکارتی آن‌ها باشد. فرض کنید نگاشت f از X به Y با ضابطه‌ی $y = f(x_1, \dots, x_n)$ باشد. در این صورت حاصل عمل f بر n مجموعه فازی $\tilde{A}_1, \dots, \tilde{A}_n$ به‌عنوان یک مجموعه فازی $\tilde{B}$ از Y به صورت زیر تعریف می‌شود.

$$\tilde{B} = f(\tilde{A}_1, \dots, \tilde{A}_n) = \left\{ (y, \tilde{B}(y)) \mid y = f(x_1, \dots, x_n), (x_1, \dots, x_n) \in X \right\} \quad (2.A)$$

که در آن

$$\tilde{B}(y) = \begin{cases} \sup_{x_1, \dots, x_n} \min \{ \tilde{A}_1(x_1), \dots, \tilde{A}_n(x_n) \}, & f^{-1}(y) \neq \emptyset \\ \circ & f^{-1}(y) = \emptyset \end{cases} \quad (3. \tilde{A})$$

۱.۲. آ برخی فواصل و مترهای فازی

فرض کنید $\tilde{A} = (m_A, l_A, r_A)_{LR}$ و $\tilde{B} = (m_B, l_B, r_B)_{LR}$ دو عدد فازی LR باشند آنگاه:

۱. فاصله یانگ و کو [۱۴۳]:

$$D_Y^{\tilde{A}, \tilde{B}} = \frac{1}{\varphi} \left((m_A - m_B)^2 + ((m_A - k_1 l_A) - (m_B - k_1 l_B))^2 + ((m_A + k_2 r_A) - (m_B + k_2 r_B))^2 \right), \quad (4. \tilde{A})$$

که در آن

$$k_1 = \int_{\circ}^1 L^{-1}(\alpha) d\alpha, \quad k_2 = \int_{\circ}^1 R^{-1}(\alpha) d\alpha.$$

و برای دو عدد فازی مثلثی متر مبتنی بر این فاصله عبارتست از

$$D_Y^{\tilde{A}, \tilde{B}} = \frac{1}{\varphi} \left((m_A - m_B)^2 + ((m_A - \frac{1}{\varphi} l_A) - (m_B - \frac{1}{\varphi} l_B))^2 + \left((m_A + \frac{1}{\varphi} r_A) - (m_B + \frac{1}{\varphi} r_B) \right)^2 \right), \quad (5. \tilde{A})$$

۲. فاصله صادقپور [۵۸]:

$$D_P^{\tilde{A}, \tilde{B}} = \begin{cases} \left[(1-q) \int_{\circ}^1 |A_{\alpha}^{-} - B_{\alpha}^{-}|^p d\alpha + q \int_{\circ}^1 |A_{\alpha}^{+} - B_{\alpha}^{+}|^p d\alpha \right]^p & p < \infty, \\ (1-q) \sup_{\circ < \alpha \leq 1} (|A_{\alpha}^{-} - B_{\alpha}^{-}|) + q \inf_{\circ < \alpha \leq 1} (|A_{\alpha}^{+} - B_{\alpha}^{+}|) & p = \infty, \end{cases}$$

که در آن $(\tilde{A})_{\alpha}^{+}$ و $(\tilde{A})_{\alpha}^{-}$ به ترتیب حد پایینی و بالایی اعداد فازی متناظر با هر α - برش است.

در حالت خاص ($p = 2$ و $q = 1/2$) و برای اعداد مثلثی، متر به صورت زیر خواهد شد:

$$D_Y^{\tilde{A}, \tilde{B}} = \frac{1}{\varphi} \left(2(m_A - m_B)^2 + ((m_A - l_A) - (m_B - l_B))^2 + ((m_A + r_A) - (m_B + r_B))^2 + ((m_A - l_A) - (m_B - l_B))(m_A - m_B) + ((m_A + r_A) - (m_B + r_B))(m_A - m_B) \right). \quad (6. \tilde{A})$$

۳. متر کلکین نما و طاهری [۸۱] :

$$D_{\mathfrak{P}}(\tilde{A}, \tilde{B}) = \frac{1}{\mathfrak{P}} \left(|m_A - m_B| + |(m_A - k_{\mathfrak{A}} l_A) - (m_B - k_{\mathfrak{A}} l_B)| \right. \\ \left. + |(m_A + k_{\mathfrak{P}} r_A) - (m_B + k_{\mathfrak{P}} r_B)| \right), \quad (7. \tilde{A})$$

و در اعداد فازی مثلثی این متر عبارتست از:

$$D_{\mathfrak{P}}(\tilde{A}, \tilde{B}) = \frac{1}{\mathfrak{P}} \left(|m_{\tilde{A}} - m_{\tilde{B}}| + |(m_{\tilde{A}} - \frac{1}{\mathfrak{P}} l_{\tilde{A}}) - (m_{\tilde{B}} - \frac{1}{\mathfrak{P}} l_{\tilde{B}})| \right. \\ \left. + |(m_{\tilde{A}} + \frac{1}{\mathfrak{P}} r_{\tilde{A}}) - (m_{\tilde{B}} + \frac{1}{\mathfrak{P}} r_{\tilde{B}})| \right). \quad (8. \tilde{A})$$

۲.۲.آ اعداد فازی

از مهم‌ترین کاربردهای اصل گسترش، تعمیم عملگرهای حسابی موجود در اعداد حقیقی به بحث فازی است که با نام اعداد فازی شناخته می‌شوند.

تعریف ۴.۲.آ. برای یک تابع $f : X \rightarrow \mathbb{R}$ و یک نقطه $y \in \mathbb{R}$ ، تعریف می‌کنیم

$$U(y) = f^{-1}([y, +\infty)) = \{x \in X : f(x) \geq y\}$$

$$L(y) = f^{-1}((-\infty, y]) = \{x \in X : f(x) \leq y\}$$

(الف) تابع f نیم‌پیوسته بالایی است اگر و فقط اگر برای هر $y \in \mathbb{R}$ داشته باشیم

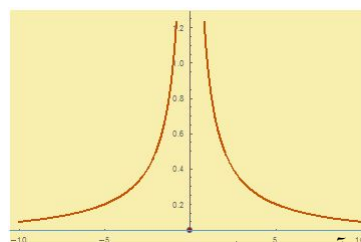
$$f\left(f^{-1}((-\infty, y))\right) \leq f(x) = y$$

(ب) تابع f نیم‌پیوسته پایینی است اگر و فقط اگر برای هر $y \in \mathbb{R}$ داشته باشیم

$$f\left(f^{-1}(y, +\infty))\right) \geq f(x) = y$$

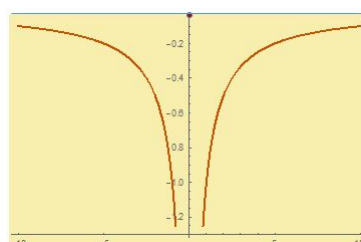
در واقع f نیم‌پیوسته بالایی (پایینی) است اگر و فقط اگر پیوسته باشد یا یک نقطه جهش بالا (پایین) داشته باشد. اگر f هم نیم‌پیوسته بالایی و هم نیم‌پیوسته پایینی باشد پیوسته است. در ادامه نمونه‌ای از تابع نیم‌پیوسته بالایی و پایینی را به همراه نمودار آن‌ها ارائه می‌نماییم.

$$f(x) = \begin{cases} \left| \frac{1}{x} \right| & x \neq 0 \\ 0 & x = 0 \end{cases}$$



شکل ۱.۰: نمودار تابع نیم‌پیوسته پایینی

$$f(x) = \begin{cases} -\left| \frac{1}{x} \right| & x \neq 0 \\ 0 & x = 0 \end{cases}$$



شکل ۲.۰: نمودار تابع نیم‌پیوسته بالایی

آلفا برش

برای $0 < \alpha \leq 1$ دلخواه، عناصری از X که درجه عضویت آن‌ها در مجموعه فازی $\tilde{A}$ حداقل به میزان α باشد، α -برش $\tilde{A}$ (مجموعه‌های تراز α وابسته به $\tilde{A}$) یا $\tilde{A}_\alpha$ گوئیم و داریم،

$$\tilde{A}_\alpha = \{x \in X, \tilde{A}(x) \geq \alpha\} \quad (9.1)$$

تعریف ۵.۲.۰. مجموعه فازی $\tilde{A}$ از $\mathbb{R}$ را یک عدد فازی گوئیم هرگاه

۱. $\tilde{A}$ نرمال و تک نمایی باشد.

۲. α -برش‌های $\tilde{A}$ ، به ازای هر $\alpha \in (0, 1]$ ، به صورت بازه‌های بسته باشد.

مجموعه اعداد فازی را با $\mathcal{F}(\mathbb{R})$ نشان می‌دهیم.

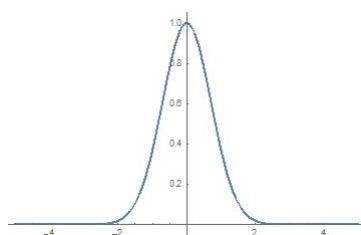
اعداد فازی دارای خواص زیر هستند،

(الف) هر عدد فازی یک مجموعه فازی محدب است.

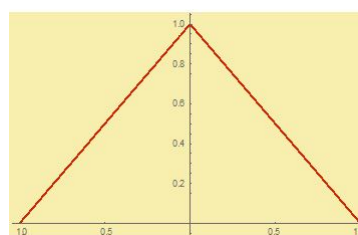
(ب) تابع عضویت هر عدد فازی، نیم‌پیوسته است.

(پ) اگر $\tilde{A}$ یک عدد فازی با $\tilde{A}(x_0) = 1$ باشد، آن‌گاه تابع عضویت $\tilde{A}$ در بازه‌ی $(-\infty, x_0)$ غیرنزولی و در بازه‌ی $(x_0, +\infty)$ غیرصعودی است.

مثال ۱.۲.۰. شکل ۳.۰ بیانگر اعداد فازی هستند که مفهوم تقریباً صفر را نمایش می‌دهند.



$$k(x) = e^{-x^2/2}, x \in \mathbb{R}. \quad (\text{ب})$$



$$g(x) = 1 - |x|, -1 \leq x \leq 1. \quad (\text{آ})$$

شکل آ.۳: نمودارهای اعداد فازی تقریباً صفر.

۳.۲.آ اعداد فازی LR

در بین اعداد فازی مختلف، اعداد فازی LR نوع خاصی از اعداد فازی هستند که علاوه بر دارا بودن ساختار ویژه، اعمال حسابی نیز بر آنها از قواعد خاصی پیروی می‌کند. به همین دلیل عمدتاً از این نوع اعداد فازی استفاده می‌گردد.

تعریف آ.۶.۲. اگر ساختار تابع عضویت عدد فازی $\tilde{M}$ به صورت زیر باشد،

$$\tilde{M}(x) = \begin{cases} L\left(\frac{m-x}{a}\right), & x < m \\ R\left(\frac{x-m}{b}\right), & x > m \end{cases}$$

آن‌گاه $\tilde{M}$ را با شرایط زیر یک عدد فازی LR می‌نامند:

(الف) L و R توابعی غیرصعودی از $\mathbb{R}^+$ به $[0, 1]$ هستند،

$$L(0) = R(0) = 1 \quad (\text{ب})$$

و آن را به صورت $\tilde{M} = (m, a, b)_{LR}$ نمایش می‌دهند، که در آن m مقدار مد (نما) و اعداد مثبت a و b به ترتیب پهنای چپ و راست $\tilde{M}$ هستند.

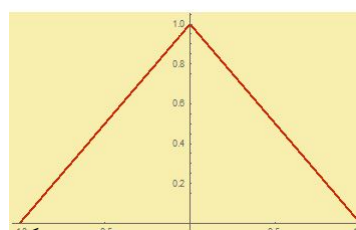
توابع زیر نمونه‌های متداولی از L (R) هستند.

$$L(x) = e^{-|x|^p} \quad L(x) = \frac{1}{1 + |x|^p} \quad L(x) = 1 - |x|^p$$

تبصره آ.۱.۲. فرض کنید $\tilde{M} = (m, a, b)_{LR}$ و $L = R$. در این صورت

(الف) $\tilde{M}$ را یک عدد فازی مثلثی گوئیم و با $\tilde{M} = (m, a, b)_T$ نمایش می‌دهیم اگر

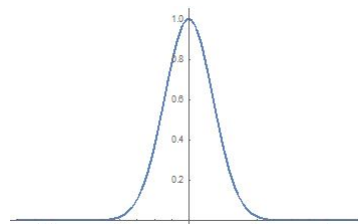
$$L(x) = \max\{0, 1 - |x|\}$$



شکل آ.۴: عدد فازی مثلثی تقریباً صفر.

ب) $\widetilde{M}$ را یک عدد فازی نرمال گوییم و با $\widetilde{M} = (m, a, b)_N$ نمایش می‌دهیم اگر

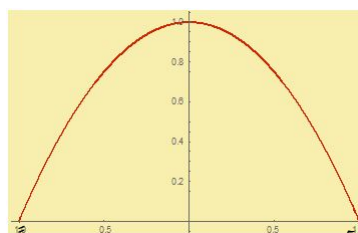
$$L(x) = e^{-x^2}$$



شکل ۵.۱: عدد فازی نرمال تقریباً صفر.

ج) $\widetilde{M}$ را یک عدد فازی سهموی گوییم و با $\widetilde{M} = (m, a, b)_P$ نمایش می‌دهیم اگر

$$L(x) = \max\{0, 1 - x^2\}$$



شکل ۶.۱: عدد فازی سهموی تقریباً صفر.

تبصره ۲.۲.آ. در اعداد فازی LR فوق‌الذکر اگر $a = b = s$ باشد، نوع متقارن آن‌ها را خواهیم داشت و به ترتیب به صورت $(m, s)_T$ ، $(m, s)_N$ و $(m, s)_P$ نمایش می‌دهیم.

۴.۲.آ عمل‌های اصلی در اعداد فازی

با توجه به اصل توسیع، چهار عمل اصلی به صورت زیر ارائه می‌گردند،

$$(\widetilde{M}_1 \oplus \widetilde{M}_2)(z) = \sup_{x,y; z=x+y} \min \{ \widetilde{M}_1(x), \widetilde{M}_2(y) \}$$

$$(\widetilde{M}_1 \ominus \widetilde{M}_2)(z) = \sup_{x,y; z=x-y} \min \{ \widetilde{M}_1(x), \widetilde{M}_2(y) \}$$

$$(\widetilde{M}_1 \otimes \widetilde{M}_2)(z) = \sup_{x,y; z=x \times y} \min \{ \widetilde{M}_1(x), \widetilde{M}_2(y) \}$$

$$(\widetilde{M}_1 \oslash \widetilde{M}_2)(z) = \sup_{x,y; z=x/y} \min \{ \widetilde{M}_1(x), \widetilde{M}_2(y) \}$$

تبصره ۳.۲.آ. اگر $\widetilde{M}_1$ و $\widetilde{M}_2$ دو عدد فازی با تابع عضویت پیوسته و پوشا بر $[0, 1]$ باشند و * یک عملگر دوتایی به‌طور پیوسته یکنوا باشد، آن‌گاه $M_1 \otimes M_2$ یک عدد فازی است که تابع عضویت آن پیوسته و بر $[0, 1]$ پوشا است.

تبصره ۴.۲.آ. اگر عملگر * دارای ویژگی جابجایی (شرکت‌پذیری) باشد، آن‌گاه عملگر $\otimes$ نیز دارای ویژگی فوق‌الذکر است.

۵.۲.آ حساب اعداد فازی LR

بواسطه عمومیت داشتن اعداد فازی LR بررسی نتایج حاصل از اثر عملگرها در آنها دارای اهمیت بسزایی است و به شرح زیر ارائه خواهد گردید:

الف) اگر $\widetilde{M} = (m, a, b)_{LR}$ و $\lambda \in \mathbb{R}$ آن گاه

$$\lambda \otimes \widetilde{M} = \begin{cases} (\lambda m, \lambda a, \lambda b), & \lambda > 0 \\ (\lambda m, -\lambda b, -\lambda a), & \lambda < 0 \end{cases} \quad (10.A)$$

ب) اگر $\widetilde{M}_1 = (m_1, a_1, b_1)_{LR}$ و $\widetilde{M}_2 = (m_2, b_2, a_2)_{LR}$ ، آن گاه داریم،

$$\widetilde{M}_1 \oplus \widetilde{M}_2 = (m_1 + m_2, a_1 + a_2, b_1 + b_2)_{LR} \quad (11.A)$$

ج) اگر $\widetilde{M}_1 = (m_1, a_1, b_1)_{LR}$ و $\widetilde{M}_2 = (m_2, b_2, a_2)_{RL}$ ، آن گاه داریم،

$$\widetilde{M}_1 \ominus \widetilde{M}_2 = (m_1 - m_2, a_1 + a_2, b_1 + b_2)_{LR} \quad (12.A)$$

د) اگر $\widetilde{M}_1 = (m_1, a_1, b_1)_{LR}$ و $\widetilde{M}_2 = (m_2, a_2, b_2)_{LR}$ ، آن گاه داریم،

$$(\widetilde{M}_1 \otimes \widetilde{M}_2) \simeq \begin{cases} (m_1 m_2, m_1 a_2 + m_2 a_1, m_1 b_2 + m_2 b_1)_{LR} & \widetilde{M}_1 \text{ و } \widetilde{M}_2 \text{ هر دو مثبت} \\ (m_1 m_2, m_1 a_2 - m_2 b_1, m_1 b_2 - m_2 a_1)_{RL} & \widetilde{M}_1 \text{ مثبت و } \widetilde{M}_2 \text{ منفی} \\ (m_1 m_2, -m_2 b_1 - m_1 b_2, -m_2 a_1 - m_1 a_2)_{RL} & \widetilde{M}_1 \text{ و } \widetilde{M}_2 \text{ هر دو منفی} \end{cases}$$

ملاحظه ۱.۲.آ. از آنجاکه تحت α -برش های مختلف نتایج حاصل از عمل ضرب، عدد فازی LR را نتیجه نمی دهد، مقدار عبارت محاسبه شده ی فوق تقریبی است.

۳.آ شاخص های نیکویی برازش

۱. متوسط خطای پیشگویی: (MPE)

$$\mathcal{P}_t = \text{MPE}_t(\widetilde{A}, \hat{\widetilde{B}}) = \frac{1}{n} \sum_{i=1}^n D_t(\widetilde{A}_i, \hat{\widetilde{B}}_i), \quad t = 1, 2, 3.$$

۲. معیار شباهت: (SIM)

$$\mathcal{S}_t = \text{SIM}_t(\widetilde{A}, \hat{\widetilde{B}}) = \frac{1}{n} \sum_{i=1}^n \frac{1}{1 + D_t(\widetilde{A}_i, \hat{\widetilde{B}}_i)}, \quad t = 1, 2, 3.$$

۴.آ نمودار تیلور

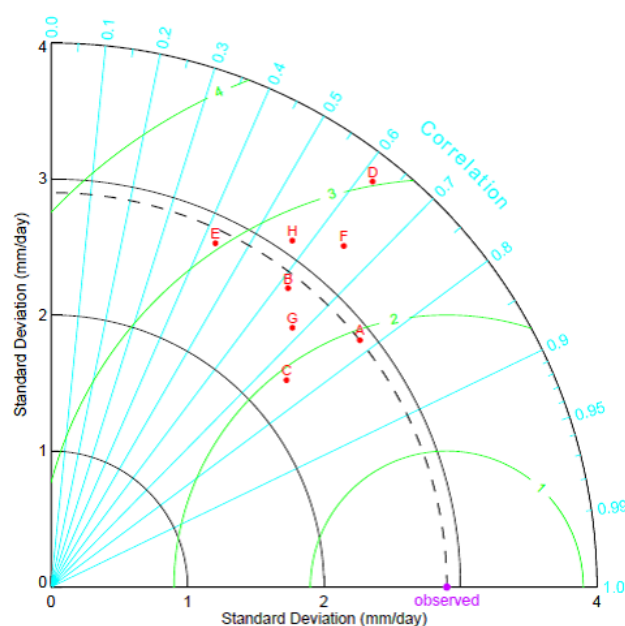
نمودار تیلور^۱ که توسط تیلور در [۱۳۳] ارائه گردید، نموداری شهودی-استنباطی است. بدین معنی که در مورد شباهت الگو بین مقدار برآورد شده و مشاهده شده به تصویر کشیده می شود. این شباهت از نظر شاخص‌های همبستگی بین برآورد و داده ها، خطای جذر میانگین توان دوم^۲ آنها و انحراف معیار آنها بررسی می شود. در واقع این نمودار برای بررسی جنبه های مختلف مدلها مناسب است (برای اطلاعات بیشتر به [۱۳۳] مراجعه شود). علاوه بر این، این نمودار مقایسه عملکرد بین مدهای مختلف را امکان پذیر کرده است. نمونه ای از این نمودار به صورت زیر ارائه می گردد.

الف. فاصله شعاعی از مبدا مختصات که نشان دهنده انحراف معیار مجموعه داده ها است.

ب. خطای RMS که نشان دهنده فاصله میان مشاهدات و برآوردهای متناسب با آنها است که در نمودار زیر با قوسهایی به مرکزیت داده مشاهده شده که روی محور افقی مشخص شده است.

ج. همبستگی میان مشاهدات و برآوردها که به صورت شعاع دایره و در زوایای مختلف ایجاد شده که مقدار آن بر روی محیط دایره نشان داده شده است.

نمونه ای از ساختار این نمودار در زیر ارائه شده است.



^۱Taylor Diagram

^۲Root mean square

پیوست ب

تعاریف و مفاهیم

ب.۱ برآوردگرهای بهبودیافته

یکی از عمده‌ترین مسائل در آمار استنباطی، ارائه یک برآوردگر مناسب برای پارامتر مورد مطالعه است. فرض کنید پارامتر مورد مطالعه در جامعه θ باشد. همچنین فرض کنید $\delta(X)$ آماره‌ای بر اساس متغیر تصادفی X بوده که به عنوان برآوردگر θ به کار می‌رود. زیان ناشی از برآورد پارامتر θ با استفاده از $\delta(X)$ را با تابع زیان^۱ $L(.,.)$ اندازه‌گیری می‌کنیم که در آن $L : (\Theta \times \chi) \rightarrow M$ و $M < \infty$ یک مقدار عددی است. در اینجا Θ فضای پارامتر و χ فضای مشاهدات است. همچنین متوسط آن تابع زیان، تابع مخاطره^۲، را با استفاده از

$$R(\theta, \delta) = E_{\theta}(L(\theta, \delta(X)))$$

محاسبه می‌کنیم.

منظور از برآوردگر بهبود یافته^۳، برآوردگری است که متوسط زیان آن کمتر از یک برآوردگر هدف (خام یا اولیه) است. لازم به ذکر است که در استنباط غیرفازی، برآوردگر اولیه معمولاً یکی از برآوردگرهای ماکزیمم درست‌نمایی، بیز، پایا و یا کمترین توان‌های دوم در مبحث رگرسیون است.

^۱Loss function

^۲Risk function

^۳Improvement estimators

از آنجایی که یکی از روش‌های کاهش مخاطره و بهبود برآوردگرها روش انقباضی^۴ است، در ادامه برآوردگر انقباضی و تعبیر هندسی نوع خاصی را برای درک موضوع به همراه یک مثال کاربردی ارائه کرده و سپس دو روش نمونه‌گیری را که منجر به افزایش کارایی به معنی کاهش مخاطره می‌شود توضیح مختصری می‌دهیم.

ب. ۱.۱. برآوردگر انقباضی

در آمار، برآوردگر انقباضی^۵ برآوردگری است که از ترکیب برآوردگر هدف با یک سری اطلاعات (که معمولاً در قالب برآوردگری برای پارامتر مورد علاقه، با ساختاری متفاوت از برآوردگر هدف است) بهبود می‌یابد. باید دقت داشت که این بهبود در جهت مقداری است که از اطلاعات دیگر بدست آمده نه در جهت برآوردگر اولیه. صورت کلی برآوردگر انقباضی عبارت از $X + g(X)$ است که در آن تابع $g(X)$ یک تابع حقیقی مقدار (اندازه‌پذیر) می‌باشد. در ادامه چند مثال از برآوردگر انقباضی آورده شده است.

۱. برآوردگر انقباضی به سمت صفر:

فرض کنید برآوردگر اولیه $\delta(X) = X$ باشد. در این صورت برآوردگر انقباضی به سمت صفر به صورت زیر خواهد بود،

$$(1 - \omega)X = X - \omega X$$

۲. برآوردگر انقباضی به سمت یک عدد ثابت m :

به‌طور مشابه برآوردگر انقباضی به سمت یک عدد ثابت m به صورت زیر خواهد بود،

$$\omega m + (1 - \omega)X = X + \omega(m - X)$$

که در آن‌ها $0 < \omega < 1$ ثابت انقباضی است و میزان انقباض را مشخص می‌کند. در ادامه نوع خاصی از برآوردگر انقباضی به نام برآوردگر انقباضی استاین^۶ را که دارای کاربرد فراوانی است مورد بحث و بررسی قرار می‌دهیم.

ب. ۲.۱. برآوردگر استاین

فرض کنید $\mathbf{X} \sim N_p(\boldsymbol{\theta}, \mathbf{I}_p)$ باشد که در آن بردار $\boldsymbol{\theta} = (\theta_1, \dots, \theta_p)^T$ نامعلوم است. استاین در سال ۱۹۵۶ نشان داد به ازای $p = 1$ ، برآوردگر $\delta(X) = X$ برای $\boldsymbol{\theta}$ مجاز است. علاوه بر آن (تحت روش بیز حدی) نشان داد که در $p = 2$ نیز مسئله برقرار است، اما زمانی که $p \geq 3$ باشد،

^۴ Shrinkage method

^۵ Shrinkage estimator

^۶ Stein shrinkage estimator

$\delta(X)$ مجاز نیست. در سال ۱۹۶۱ استاین به همراه شاگردش جیمز برآوردگری به صورت زیر ارائه نمودند.

$$\delta^{JS}(X) = \left(1 - \frac{a}{\|X\|^2}\right) X, \quad 0 < a < 2(p-2) \quad (۱.ب)$$

$$= X - \frac{a}{\|X\|^2} X, \quad 0 < a < 2(p-2) \quad (۲.ب)$$

$$= X - g(a)X \quad (۳.ب)$$

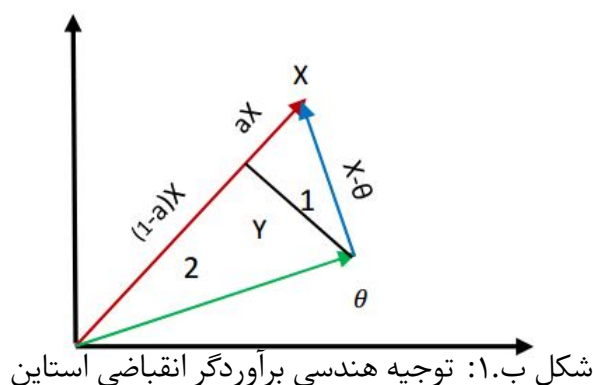
که در آن $g(a) = \frac{a}{\|X\|^2}$ ضریب انقباض نامیده می‌شود و با افزایش مقدار $g(a)$ (نزدیک شدن به ۱) انقباض به سمت صفر را خواهیم داشت. در بخش ۳.۲.۲ نشان می‌دهیم مخاطره برآوردگری با این ساختار از برآوردگر معمولی $\delta(X) = X$ برای $p \geq 3$ کمتر است. در ادامه برای درک بهتر برتری برآوردگر انقباضی استاین تعبیر هندسی آن آورده می‌شود.

۳.۱.ب تعبیر هندسی برآوردگر نوع استاین

فرض کنید $X = (X_1, X_2, \dots, X_p)^T$ یک بردار p بعدی با مولفه‌های مستقل یا ناهمبسته و میانگین θ باشد و $E(X) = \theta$ و $V(X_i) = \sigma^2$. واضح است که $E((X - \theta)^T \theta) = 0$ زیرا همچنین فرض کنید

$$E(X^T \theta - \theta^T \theta) = (E(X))^T \theta - \theta^T \theta = \theta^T \theta - \theta^T \theta = 0$$

بنابراین انتظار داریم $X - \theta$ بر θ عمود باشد. (مطابق شکل ۱.ب)



چون $E(\|X\|^2) = \text{Var}(\|X\|) + (E(X))^T(E(X)) = p\sigma^2 + \theta^T \theta = p\sigma^2 + \|\theta\|^2$ برآورد $\|\theta\|^2$ بیش برآوردی به اندازه $p\sigma^2$ داریم. در نتیجه می‌توان انتظار داشت تصویر θ روی X برآوردگر مناسب‌تری باشد. اما باید توجه داشت که جنس تصویر θ از نوع θ می‌باشد. لذا برآوردگر معقولی نیست ولی می‌توان آن را برآورد کرد. حال اگر این تصویر را با $(1-a)X$ نشان دهیم، هدف مورد نظر برآورد a می‌باشد. یک راه برآورد a استفاده از قضیه‌ی فیثاغورث

می باشد. از آن جایی که $E(\|X\|^2)$ تخمینی برای $\|X\|^2$ است. انتظار داریم این دو از هم زیاد فاصله نداشته باشد. اگر $\|X\|^2 = p\sigma^2 + \|\theta\|^2$ ، به همین ترتیب $\|X - \theta\|^2 = p\sigma^2$. با استفاده از قضیه فیثاغورث در مثلث ۱ داریم،

$$\begin{aligned}\|Y\|^2 &= \|X - \theta\|^2 - \hat{a}^2 \|X\|^2 \\ &= p\sigma^2 - \hat{a}^2 \|X\|^2\end{aligned}\quad (\text{ب.۴})$$

در مثلث ۲ داریم،

$$\begin{aligned}\|Y\|^2 &= \|\theta\|^2 - (1 - \hat{a})^2 \|X\|^2 \\ &= \|X\|^2 - p\sigma^2 - (1 - \hat{a})^2 \|X\|^2\end{aligned}\quad (\text{ب.۵})$$

با تلفیق روابط (ب.۴) و (ب.۵) داریم،

$$\begin{aligned}p\sigma^2 - \hat{a}^2 \|X\|^2 &= \|X\|^2 - p\sigma^2 - (1 - \hat{a})^2 \|X\|^2 \\ \Rightarrow \hat{a} &= \frac{p\sigma^2}{\|X\|^2} \quad (1 - \hat{a})X = \left(1 - \frac{p\sigma^2}{\|X\|^2}\right)X\end{aligned}$$

در اینجا به چند نکته مهم می توان اشاره کرد،

۱. روش فوق به نرمال بودن توزیع X و اینکه θ پارامتر مکان باشد بستگی ندارد.
 ۲. روش فوق برای غیر مجاز بودن X و تعیین مقدار p مناسب است.
 ۳. روش فوق امکان بدست آوردن برآوردگر بهتری نسبت به برآوردگر ناریب را با منقبض کردن برآوردگر X به سمت مرکز فراهم می کند.
- در ادامه به مثال زیر از صالح (۲۰۰۶) توجه کنید، که رفتار برآوردگر انقباضی را به طور شهودی نشان می دهد.

مثال ۱.۱.۱.ب. مسابقه بیس بال را در نظر بگیرید. این ورزش گروهی با چوب و توپ توسط ۹ بازیکن انجام می شود. برای فهم بیشتر بهتر است با برخی اصطلاحات این ورزش آشنا شویم،

- **بتر رانر:** به فردی اطلاق می شود که بعد از ضربه زدن به توپ در حال دویدن به بیس اول باشد و تا زمانی که به بیس اول نرسیده بتر رانر است و بعد از اینکه به بیس اول رسید در اصطلاح رانر گفته می شود. تمام تلاش بتر رانر باید این باشد که خود را به بیس اول برساند زیرا در غیر این صورت او اوت شده و به رانر برای کسب امتیاز تبدیل نمی شود.

- **بتر باکس:** مکانی است در دو طرف بیس خانه به صورت مستطیل شکل و به طول ۱/۵۲ و عرض ۱/۰۲ که بتر در آنجا می ایستد تا توپ پرتاب شده توسط پیچر را بزند. بتر برای ضربه زدن نباید پایش را از منطقه بیرون بگذارد و تنها با اجازه داور می تواند از این محدوده بیرون برود.

● بتینگ: عملی است که یک بتر به عنوان بازیکن تیم مهاجم به تنهایی در مقابل مدافعان انجام می‌دهد و به معنی ضربه زدن توپ پرتاب شده توسط پیچر است و با این کار به حمله و مبارزه با تیم مهاجم برای کسب امتیاز خودش و یارانش می‌پردازد.

● میانگین بتینگ: میانگین بتینگ در آمار بیس بال برای شناخت عملکرد بتر بسیار موثر است. این آمار را هنری چاودیک در ورزش کریکت اختراع کرد.

به عنوان مثال سه بازیکن یا بیشتر را در یک بازی بیس بال در نظر بگیرید. علاقه‌مندیم بتینگ اوریج را برای هر یک از بازیکنان پیش‌بینی کنیم. آماردان‌هایی که از روش انقباضی استفاده می‌کنند انتظار دارند تا ضربه‌های آینده بازیکنان را دقیق‌تر از مقادیر از مقادیر پیش‌بینی شده براساس بتینگ اوریج جداگانه‌ی هر بازیکن پیش‌بینی کنند. روش استاین بر پایه‌ی داده‌های بیس بال مشهور افرون و موریس را در زیر نشان می‌دهیم، جدول ۱.۱ میانگین بتینگ ۴۵ ضربه‌ی اول هر ۱۸ بازیکن را نشان می‌دهد،

جدول ۱.۱: نتایج بازی بیس بال مربوط به مثال ۱.۱.۱.

ردیف	میانگین بعد از ۴۵ ضربه‌ی اول ($\bar{X}_i$)	پیش‌گویی میانگین به روش انقباضی ($\bar{X}_i^{JS}$)	ردیف	میانگین بعد از ۴۵ ضربه‌ی اول ($\bar{X}_i$)	پیش‌گویی میانگین به روش انقباضی ($\bar{X}_i^{JS}$)
۱	۰/۴۰۰	۰/۲۹۴	۱۰	۰/۲۴۴	۰/۲۶۱
۲	۰/۴۷۸	۰/۲۸۹	۱۱	۰/۲۲۲	۰/۲۵۶
۳	۰/۳۵۶	۰/۲۴۵	۱۲	۰/۲۲۲	۰/۲۵۶
۴	۰/۳۳۳	۰/۲۸۰	۱۳	۰/۲۲۲	۰/۲۵۶
۵	۰/۳۱۱	۰/۲۷۵	۱۴	۰/۲۲۲	۰/۲۵۶
۶	۰/۳۱۱	۰/۲۷۵	۱۵	۰/۲۲۲	۰/۲۵۶
۷	۰/۲۸۹	۰/۲۷۰	۱۶	۰/۲۰۰	۰/۲۵۲
۸	۰/۲۶۷	۰/۲۶۶	۱۷	۰/۱۷۸	۰/۲۴۱
۹	۰/۲۴۴	۰/۲۶۱	۱۸	۰/۱۵۶	۰/۲۴۳

نخستین مرحله در روش استاین محاسبه‌ی میانگین کل تمام میانگین‌هاست. که در اینجا برابر است با، $\bar{X}_T = \frac{\bar{X}_1 + \bar{X}_2 + \dots + \bar{X}_{18}}{18} = ۰/۲۶۵۴$.

هدف، انقباض ۱۸ میانگین به سمت میانگین کل است که منطقی به نظر می‌رسد. مرحله‌ی بعد یافتن توان دوم فاصله‌ی وزنی بین ۱۸ میانگین جدول و میانگین کل است

که در این مثال ۱۹/۰۴۵ است. به دنبال این مرحله ثابت انقباضی

$$c = \left(1 - \frac{2 - \text{درجه آزادی}}{\text{مربع فاصله وزنی}} \right) = \left(1 - \frac{15}{19/0.45} \right) = 0.212$$

را بدست می‌آوریم. عدد ۱۵ همان درجه‌ی آزادی فاصله‌ی وزنی منهای ۲ است. دقت کنید ثابت انقباضی را از آن جهت ثابت گویند که برای هر ۱۸ بازیکن ثابت است و در محاسبه‌ی آن باید این نکته را مد نظر داشت که تنها میانگین بتینگ ۱۷ بازیکن دیگر بر بازیکن مد نظر دخالت دارد و لذا درجه‌ی آزادی برابر عدد ۱۵ خواهد بود. پیش‌بینی میانگین به روش انقباضی جیمز-استاین از رابطه‌ی زیر بدست می‌آید،

$$\bar{X}_i^{JS} = \bar{X}_T + c(\bar{X}_i - \bar{X}_T), \quad i = 1, 2, \dots, 18$$

به عنوان مثال برای بازیکن‌های اول و هجدهم به ترتیب داریم،

$$\bar{X}_1^{JS} = \bar{X}_T + c(\bar{X}_1 - \bar{X}_T) = 0.265 + 0.212(0.400 - 0.265) = 0.294$$

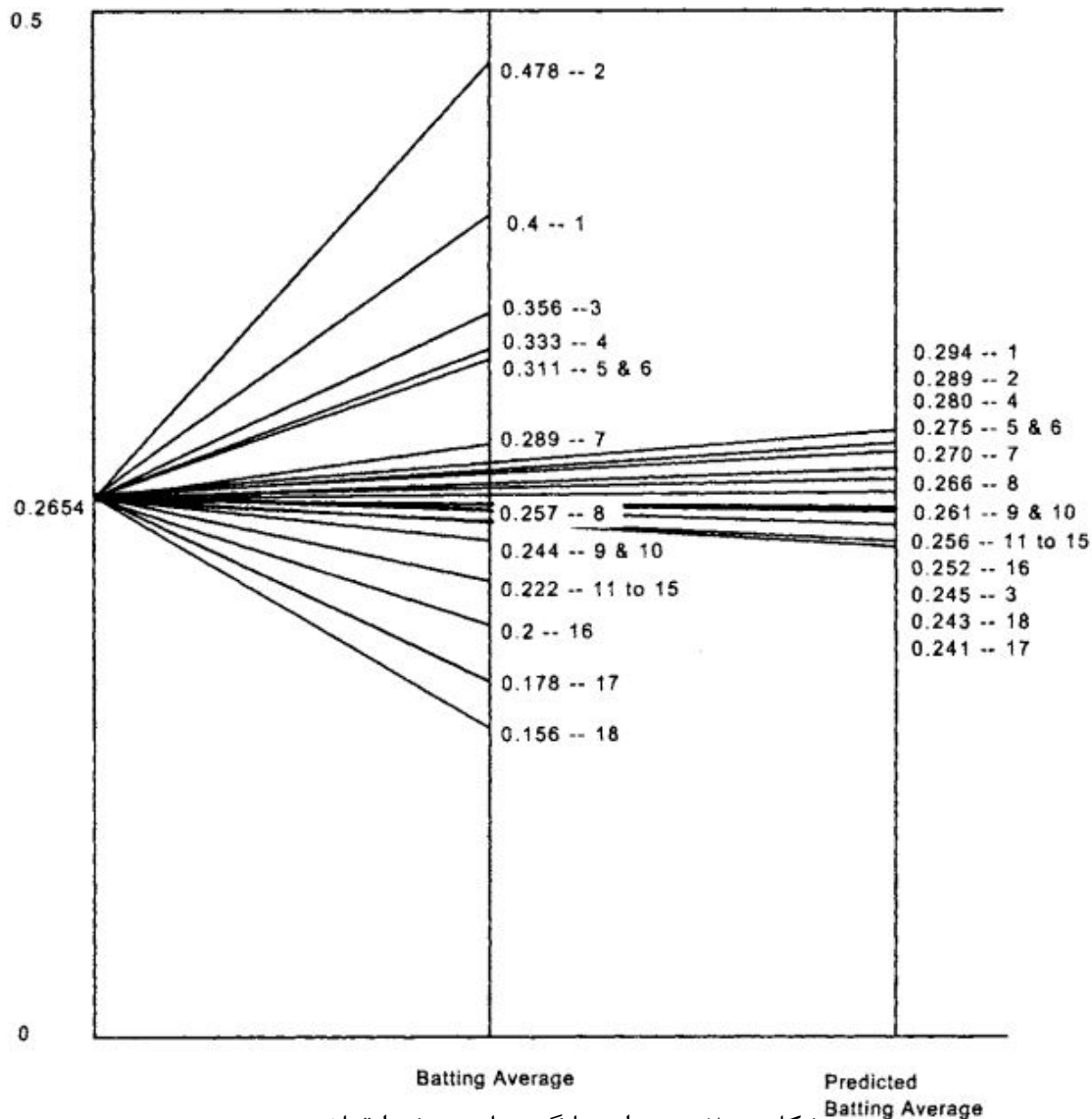
$$\bar{X}_{18}^{JS} = \bar{X}_T + c(\bar{X}_{18} - \bar{X}_T) = 0.265 + 0.212(0.156 - 0.265) = 0.243$$

در حالی که میانگین واقعی این دو بازیکن به ترتیب ۰/۳۴۶ و ۰/۲۰۰ می‌باشد.

از شکل ۲.ب مشاهده می‌شود که میانگین‌های بدست آمده به روش انقباضی در مقایسه با میانگین بتینگ معمولی بیشتر در اطراف میانگین کل جمع شده‌اند. به عبارت دیگر روش استاین برآوردهای اولیه را به سمت میانگین کل هدایت می‌کند.

ب.۴.۱ بررسی رفتار برآوردهای انقباضی در توزیع نرمال

در این بخش نشان می‌دهیم که برآوردهای انقباضی مخاطره کمتری نسبت به برآوردهای معمولی در $p \geq 3$ دارا می‌باشد.



شکل ب.۲: نمودار میانگین‌ها به روش انقباضی

قضیه ب.۱.۱. اگر $\mathbf{X} \sim N_p(\boldsymbol{\theta}, \mathbf{I}_p)$ آن‌گاه تحت تابع زیان توان دوم خطا، برآوردهای $\delta^{JS}(\mathbf{X})$ بر برآوردهای $\mathbf{X}$ تحت شرایط $p \geq 3$ و $0 < a < 2(p-2)$ برتری دارد همچنین برآوردهای $\delta^{JS}(\mathbf{X}) = \left(1 - \frac{p-2}{\|\mathbf{X}\|^2}\right) \mathbf{X}$ نسبت به هر برآوردهای دیگری در این کلاس به‌طور یکنواخت دارای کمترین مخاطره است.

برهان. طبق تعریف برآوردهای انقباضی استاین داریم،

$$\begin{aligned} R(\boldsymbol{\theta}, \delta^{JS} \mathbf{X}) &= E \left[\|\delta^{JS}(\mathbf{X}) - \boldsymbol{\theta}\|^2 \right] \\ &= E \left[\left\| \left(1 - \frac{a}{\|\mathbf{X}\|^2} \right) \mathbf{X} - \boldsymbol{\theta} \right\|^2 \right] \\ &= E \left[\|\mathbf{X} - \boldsymbol{\theta}\|^2 \right] + a^2 E \left[\frac{1}{\|\mathbf{X}\|^2} \right] - 2a E \left[\frac{(\mathbf{X} - \boldsymbol{\theta})^T \mathbf{X}}{\|\mathbf{X}\|^2} \right] \\ &= p + a^2 E \left[\frac{1}{\|\mathbf{X}\|^2} \right] - 2a \sum_{i=1}^p E \left[\frac{(X_i - \theta_i) X_i}{\|\mathbf{X}\|^2} \right] \end{aligned}$$

با بکار بردن لم استاین [۱۲۴] برای $h_i(\mathbf{X}) = \frac{X_i}{\|\mathbf{X}\|^2}$ ، باتوجه به اینکه $\frac{\partial \|\mathbf{X}\|^2}{\partial X_i} = \frac{\partial \sum X_i}{\partial X_i} = 2X_i$ داریم،

$$\frac{\partial h_i(\mathbf{X})}{\partial X_i} = \frac{\|\mathbf{X}\|^2 - X_i(2X_i)}{\|\mathbf{X}\|^4}$$

لذا می‌توان نتیجه گرفت

$$\begin{aligned} R(\boldsymbol{\theta}, \delta^{JS}(\mathbf{X})) &= p + a^2 E \left[\frac{1}{\|\mathbf{X}\|^2} \right] - 2a \sum_{i=1}^p E \left[\frac{\partial h_i(\mathbf{X})}{\partial X_i} \right] \\ &= p + a^2 E \left[\frac{1}{\|\mathbf{X}\|^2} \right] - 2a \sum_{i=1}^p E \left[\frac{\|\mathbf{X}\|^2 - 2X_i^2}{\|\mathbf{X}\|^4} \right] \\ &= p + a^2 E \left[\frac{1}{\|\mathbf{X}\|^2} \right] - 2a E \left[\frac{\sum_{i=1}^p \|\mathbf{X}\|^2 - 2 \sum_{i=1}^p X_i^2}{\|\mathbf{X}\|^4} \right] \\ &= p + a^2 E \left[\frac{1}{\|\mathbf{X}\|^2} \right] - 2a E \left[\frac{p\|\mathbf{X}\|^2 - 2\|\mathbf{X}\|^2}{\|\mathbf{X}\|^4} \right] \\ &= p + E \left[\frac{a^2 - 2a(p-2)}{\|\mathbf{X}\|^2} \right] \\ &= E \left[p + \frac{a^2 - 2a(p-2)}{\|\mathbf{X}\|^2} \right] \end{aligned}$$

از طرفی $R(\boldsymbol{\theta}, \mathbf{X}) = p$ بنابراین

$$\begin{aligned} R(\boldsymbol{\theta}, \delta^{JS}(\mathbf{X})) - R(\boldsymbol{\theta}, \mathbf{X}) &\leq 0 \iff E \left[\frac{a^2 - 2a(p-2)}{\|\mathbf{X}\|^2} \right] \leq 0 \\ &\iff a^2 - 2a(p-2) \leq 0 \\ &\iff a(a - 2(p-2)) \leq 0 \\ &\iff 0 < a < 2(p-2) \end{aligned}$$

زیرا

	0 < 2(p-2)		
	+	-	+

زمانی که $a = p - 2$ است مخاطره‌ی $\delta^{JS}(\mathbf{X})$ کمترین مقدار است و داریم

$$R(\boldsymbol{\theta}, \delta^{JS}(\mathbf{X})) = p - E \left[\frac{(p-2)^2}{\|\mathbf{X}\|^2} \right]$$

□

اکنون در ادامه دیدگاه دوم را مورد بررسی قرار داده و روش‌های آن را بیان می‌نماییم.

ب.۲ روش‌های باز نمونه‌گیری

در بسیاری از مسائل واقعی ما با دو مشکل مواجه هستیم، مشکل اول ندانستن توزیع مشاهدات و مشکل دوم کم بودن تعداد مشاهدات است. از این رو مجبور به بررسی و تحلیل مشاهدات از روش‌های ناپارامتریک هستیم. در این بین روش‌هایی با نام «روش‌های باز نمونه‌گیری» وجود دارد که ما را در افزایش نمونه مشاهده شده، تخمین اریبی برآورد و خطای استاندارد و دیگر مسائل استنباطی کمک خواهد کرد.

ب.۲.۱ روش جک‌نایف

در سال ۱۹۵۶ کونویلی^۷ روش باز نمونه‌گیری بنام جک‌نایف^۸ را ارائه و از آن برای برآورد اریبی پارامتر مورد مطالعه استفاده کرد. سپس در سال ۱۹۵۸ توکی^۹ آن را برای برآورد خطای استاندارد بکار برد. در این روش با حذف هر بار یک نمونه، براساس $n - 1$ نمونه مانده کار می‌کند. بدین ترتیب طی n بار تکرار، n نمونه $n - 1$ تایی می‌سازیم و برآوردها را از آن بدست می‌آوریم. اگر $\hat{\theta}_n^*$ برآورد جک‌نایف برای پارامتر θ باشد داریم،

$$\hat{\theta}_n^* = \frac{1}{n} \sum_{i=1}^n \hat{\theta}_i$$

که در آن $\hat{\theta}_i$ برآورد حاصل از نمونه i ام است. حال اگر $\hat{\theta} = \frac{1}{n} \sum_{i=1}^n (x_i - \bar{x})^2$ برآورد واریانس X فرض شود و σ_X^2 واریانس واقعی X باشد داریم،

$$\text{اریبی}(\hat{\theta}) = \frac{-\sigma_X^2}{n} \quad \text{اریبی}(\hat{\theta}_i) = \frac{-\sigma_X^2}{n-1}$$

پس اختلاف بین این دو اریبی برابر است با، $\frac{-\sigma_X^2}{n(n-1)}$. که نشان از مناسب بودن این روش در برآورد $\hat{\theta}$ دارد. همچنین خطای استاندارد این روش عبارت است از،

$$\hat{s}_{\text{jack}} = \sqrt{\frac{n-1}{n} \sum_{i=1}^n (\hat{\theta}_i - \hat{\theta}_n^*)^2} = \left(\frac{n-1}{n}\right)^{\frac{1}{2}} \hat{s}_e(\bar{X})$$

البته این روش زمانی که پارامتر θ هموار نباشد دچار اشکال خواهد شد که برای رفع آن می‌توان به [۴۸] اشاره کرد.

^۷Quenouille

^۸Jackknife

^۹Tukey

ب.۲.۲ روش بوت‌استرپ

این روش در سال ۱۹۷۹ توسط افرون^{۱۰} [۴۷] ابداع گردید و در سال‌های بعد آن را توسعه داد. روش‌های بوت‌استرپ^{۱۱} یک کلاسی از روش‌های ناپارامتریک هستند که به برآورد توزیع یک جامعه با استفاده از باز نمونه‌گیری می‌پردازد. این روش می‌تواند در دو حالت با جایگذاری و بدون جایگذاری نمونه‌ها انجام بگیرد. کاربرد این روش در رگرسیون به دو شکل انجام می‌پذیرد. در شکل اول روی ماتریس داده‌های ورودی و جایگزینی هر سطر از آن کار می‌شود و در شکل دوم جابجایی خطاهای حاصل از برآورد اولیه تحت نمونه‌های اولیه را خواهیم داشت. تحت این روش می‌توان نمونه‌هایی از توزیع تجربی مشاهدات تولید نماید. همچنین میزان اریبی و خطای استاندارد این روش نسبت به جک‌نایف، بواسطه نامحدود بودن تعداد تکرارها در نمونه‌ها کمتر خواهد شد. اما مزیت روش جک‌نایف آن است که یک تقریب برای خطای استاندارد ارائه می‌دهد و بازه اطمینان کوچکتری نسبت به بوت‌استرپ دارد.

ب.۳.۲ روش جک‌نایف بعد از بوت‌استرپ (JB)

در روش بوت‌استرپ بواسطه تصادفی بودن تعداد نمونه‌های تکراری حاصله، خطای استاندارد و اریبی بدست آمده یک متغیر تصادفی خواهد بود. از این‌رو افرون و تیبشیرانی^{۱۲} در سال ۱۹۹۴ در [۴۸] برای بدست آوردن یک مقدار تقریبی برای خطای استاندارد از روش جک‌نایف کمک گرفته و جک‌نایف بعد از بوت‌استرپ^{۱۳} را ارائه کردند. علاوه بر آن، آن‌ها در [۴۸] نشان دادند که این روش در کاهش خطای استاندارد نسبت به بوت‌استرپ موثر است. در برآورد پارامترها برای نمایش مزیت و مقایسه روش‌های پیشنهادی با روش‌های دیگر در مدل‌های رگرسیون فازی، از مترها و شاخص‌هایی که بر مبنای فواصل بین پاسخ‌های مشاهده شده و برآورد شده ارائه شده‌اند و همچنین نمودارهایی خاص استفاده خواهیم کرد. لذا در این ضمیمه به معرفی تعدادی از این ابزارها که در رساله بکار رفته است خواهیم پرداخت.

^{۱۰}Efron

^{۱۱}Bootstrap

^{۱۲}Tibshirani

^{۱۳}Jackknif-after-Bootstrap

مراجع

- [۱] داوودی، م. ج.، (۱۳۹۴)، پایان‌نامه ارشد: ”رگرسیون فازی ناپارامتریک با استفاده از موجکها“، دانشکده علوم ریاضی و کامپیوتر، دانشگاه صنعتی امیر کبیر.
- [۲] طاهری سید محمود، ماشین‌چی ماشالله، (۱۳۸۷)، ”مقدمه‌ای بر احتمال و آمار فازی“، ویرایش اول، شهید باهنر کرمان.
- [۳] طاهری سید محمود، (۱۳۹۶)، ”رویکردهای ابتکاری در رگرسیون فازی“، اندیشه آماری، سال هجدهم، شماره دوم، صفحه ۴۳-۵۲.
- [۴] عرب‌پور ع. و مرادی ف.، (۱۳۹۳)، ”استفاده از روش بوت‌استرپ در رگرسیون مربعات فازی“، اندیشه آماری، سال پانزدهم، شماره اول، صفحه ۵۶-۶۸.
- [۵] ماشین‌چی، م. (۱۳۹۷) ”رگرسیون با استفاده از پایگاه اطلاعاتی مشکک“، اندیشه آماری، سال پنجم، شماره دوم، صفحه ۱۳-۱۹.
- [۶] مقدس م.، (۱۳۹۲)، پایان‌نامه ارشد: ”مقایسه چند روش بهینه‌سازی در رگرسیون خطی فازی“، دانشکده علوم ریاضی، دانشگاه فردوسی مشهد.
- [۷] میرزایی، ی. و ارقامی، ن. (۱۳۸۶) ”رگرسیون فازی: مروری بر چند رویکرد“، اندیشه آماری، سال دوازدهم، شماره اول، صفحه ۳۵-۴۷.
- [8] Akbari M. G., Mohammadalizadeh R., and Rezaei M. (2012). “Bootstrap statistical inference about the regression coefficients based on fuzzy data”, **International Journal of Fuzzy Systems**, 14(4), pp. 549-556.
- [9] Akbari M. G., and Hesamian G. (2019) “Elastic net oriented to fuzzy semiparametric regression model with fuzzy explanatory variables and fuzzy responses” **IEEE Transactions on Fuzzy Systems**, 27(12), pp. 2433-2442
- [10] Akbari, M. G. and Hesamian, G. (2019), “Elastic net oriented to fuzzy semiparametric regression model with fuzzy explanatory variables and fuzzy responses” **IEEE Transactions on Fuzzy Systems**, 27(12), p. 2433-2442.

-
- [11] Algur, S. P. and Biradar, J. G. (2017), "Cooks Distance and Mahanabolis Distance Outlier Detection Methods to identify Review Spam" **International Journal of Engineering and Computer Science**, 6(6), pp. 21638-21649.
- [12] Angers, J-F. and Delampady, M. (2008), "Fuzzy sets in nonparametric Bayes regression" **Pushing the Limits of Contemporary Statistics: Contributions in Honor of Jayanta K. Ghosh**, 3, pp. 89-104.
- [13] Arabpour A.R., Tata M. (2008), "Estimating the parameters of a fuzzy linear regression model" **Iranian Journal of Fuzzy Systems**, 5(2), pp. 1-19.
- [14] Arefi M. (2016), "Clustering regression based on interval-valued fuzzy outputs and interval-valued fuzzy parameters", **International Journal of Fuzzy Systems**, 30, pp. 1331-1337.
- [15] Asai H., Tanaka S. and Uegima K. (1982), "Linear regression analysis with fuzzy model" **IEEE Transaction Systems Man and Cybermatics** 12(6), pp. 903-907.
- [16] Baranchik A.J. (1970), "A family of minimax estimators of the mean of a multivariate normal distribution" **The Annals of Mathematical Statistics**, pp. 642-645.
- [17] Bargiela, A. , Pedrycz, W. ,and Nakashima, T. (2007), "Multiple regression with fuzzy data" **Fuzzy Sets and Systems**, 158(19), pp. 2169-2188.
- [18] Belsley, D. A. , Kuh, E. ,and Welsch, R. E. (1980), "Regression Diagnostic" **John Wiley**, New York.
- [19] Buckley, J. J. and Eslami, E. (2004), "Uncertain Probability II: the continuous case" **Soft Computing**, 8, pp. 193-199.
- [20] Buckley, J. J. (2005a), "Fuzzy statistics:hypothesis testing" **Soft Computing**, 9, pp. 512-518.
- [21] Buckley, J. J. (2005b), "Fuzzy statistics:regression and prediction" **Soft Computing**, 9, pp. 769-775.
- [22] Buckley, J. J. (2006), "Fuzzy probability and statistics" **Heidelberg:Springer**
- [23] Celmiņš A. (1987), "Least squares model fitting to fuzzy vector data" **Fuzzy Sets and Systems**, 22(3), pp. 245-269.

- [24] Černý, M., and Hladík, M. (2018). "Possibilistic linear regression with fuzzy data: Tolerance approach with prior information" **Fuzzy Sets and Systems**, 340, pp. 127-144.
- [25] Chachi J. and Roozbeh, M. (2017) "A fuzzy robust regression approach applied to bedload transport data" **Communications in Statistics-Simulation and Computation**, 46(3), pp. .1703-1714
- [26] Chachi J., Taheri S.M. and Arghami N.R. (2014), "A hybrid fuzzy regression model and its application in hydrology engineering", **Applied Soft Computing**, 25, 149-158.
- [27] Chachi J., Taheri S.M. and Rezaee Pazhand H. (2016), "Suspended load estimation using L1-fuzzy regression, L2-fuzzy regression and MARS-fuzzy regression models", **Hydrological Sciences Journal**, 61, pp. 1489- 1502.
- [28] Chachi, J., Taheri, S. M., Fattahi, S., and Hosseini Ravandi, S. A. (2016), "Two Robust Fuzzy Regression Models and Their Applications in Predicting Imperfections of Cotton Yarn" **Journal of Textiles and Polymers**, 4(2), pp. 60-68.
- [29] Chakraborti C. and Chakraborti D. (2008), "Fuzzy linear and polynomial regression modelling of If-Then fuzzy rule-base", **International Journal of Fuzzy Systems**, 16, pp. 659-671.
- [30] Chang P.T. and Lee E.S. (1994), "Fuzzy least absolute deviations regression based on the ranking of fuzzy numbers" **In Fuzzy Systems**, IEEE World Congress on Computational Intelligence, Proceedings of the Third IEEE Conference on pp. 1365-1369.
- [31] Chen S.P. and Dang J.F. (2008), "A variable spread fuzzy linear regression model with higher explanatory power and forecasting accuracy" **Information Sciences**, 178(20), pp. 3973-3988.
- [32] Chen L. H., Hsueh C. C. and Chang C. J. (2013) "A two-stage approach for formulating fuzzy regression models" **Knowledge-based systems**, 52, pp. .302-310
- [33] Ciavelino E. and Calcagni A. (2016), "A generalized maximum entropy (GME) estimation approach to fuzzy regression model", **Applied Soft Computing**, 38, pp. 51-63.

- [34] Choi S.H. and Buckley J.J. (2008), “Fuzzy regression using least absolute deviation estimators” **Soft Computing**, 12(3), pp. 257-263.
- [35] Cook, R. D. (1977), “Detection of Influential Observation in Linear Regression” **Technometrics**, 19(1), pp. 15-18.
- [36] Coppi R. D’Urso P. Giordani P. Santoro A.(2006), “Least square estimation of a linear regression model with LR fuzzy response” **Computational Statistics & Analysis**, 51(1), pp. 267-286.
- [37] Dang, W. , Liu, H. , Xu, J. , Zhao, H. ,and Song, Y. (2020), “An Improved Quantum-Inspired Differential Evolution Algorithm for Deep Belief Network” **IEEE Transactions on Instrumentation and Measurement**, pp. 1-1.
- [38] Dang, W. , Xu, J. and Zhao, H. (2019), “An improved ant colony optimization algorithm based on hybrid strategies for scheduling problem” **IEEE access**, 7, pp. 20281-20292.
- [39] Dariane, A. B., Azimi, S. (2018), “Stream flow forecasting by combining neural networks and fuzzy models using advanced methods of input variable selection” **Journal of Hydro-informatics**, 20(2), pp. 520-532.
- [40] Diamond P. (1987), “Fuzzy least squares” **Information Sciences**, 46(3), pp. 141-157.
- [41] Dubois D. J. (1980), “**Fuzzy Sets and Systems: Theory and Applications**”, Academic Press, New York.
- [42] D’Urso P. (2003), “Linear regression analysis for fuzzy/crisp input and fuzzy/crisp output data” **Computational Statistics & Data Analysis**, 42(1-2), pp. 47-72.
- [43] D’Urso P. and Gastaldi T. (2000), “A least-squares approach to fuzzy linear regression analysis” **Computational Statistics and Data Analysis**, 34(4), pp. 427-440.
- [44] D’Urso P. and Gastaldi T. (2002), “An “order wise” polynomial regression procedure for fuzzy data” **Fuzzy sets and systems**, 130(1), pp. 1-19.
- [45] D’Urso P. Massari R. Santoro A. (2011), “Robust fuzzy regression analysis” **Information Sciences**, 181(19), pp. 4154-4174.

- [46] D'Urso P. Massari R. (2013), "Weighted least squares and least median squares estimation for the fuzzy linear regression analysis" **Metron**, 71(3), pp. 279-306.
- [47] Efron B. (1979), "Bootstrap methods: another look at the jackknife," **Annals of Statistics**, 7(1), pp. 1-26.
- [48] Efron B. and Tibshirani R.J. (1994), **An introduction to the bootstrap**, Chapman & Hall, CRC, Boca Raton, SL.
- [49] Falsafian, A. , Taheri S. M. , and Mashinchi M. (2008), "Fuzzy Estimation of Parameters in Statistics Models" **Int. J. Camp. Math. Sci**, pp. 79-85.
- [50] Falsafian, A. and Taheri S. M. (2011), "On Buckley's approach to fuzzy estimation" **Soft Computing**, 15, pp. 345-349.
- [51] Fan, J. and Li, R. (2001), "Variable selection via nonconcave penalized likelihood and its oracle properties" **Journal of the American statistical Association**, 96(456) pp. 1348-1360.
- [52] Farnoosh R., Ghasemian J. and Soleymani Fard A. (2012), "A modification on ridge estimation for fuzzy nonparametric regression" **Iranian Journal of Fuzzy Systems**, 9, pp. 75-88.
- [53] Ferraro M. B., Coppi R., and González-Rodríguez G. (2013), "Bootstrap confidence intervals for the parameters of a linear regression model with fuzzy random variables In Towards Advanced Data Analysis by Combining" **Soft Computing and Statistics**, pp. 33-42.
- [54] Fox J. and Monette G. (1992), "Generalized Collinearity Diagnostics" **Journal of the American Statistical Association**, 87(417) pp. 178-183.
- [55] Frank L. E. and Fridman J. H. (1993), "A statistical view of some chemometrics regression tools" **Technometrics**, 35(2) pp. 109-135.
- [56] Fu, W. J. (1998), "Penalized regressions: the bridge versus the lasso" **Journal of computational and graphical statistics**, 7(3) pp. 397-416.
- [57] Gao, Y., and Lu, Q. (2018), "A fuzzy logistic regression model based on the least squares estimation" **Computational and Applied Mathematics**, 37(3), pp. 3562-3579.

- [58] Gildeh B. S. and Gien D. E. N. I. S. (2002), “A goodness of fit index to reliability analysis in fuzzy model” **Advances in Intelligent Systems, Fuzzy Systems, Evolutionary Computation**, pp. 78-83.
- [59] Hao, P. and Chiang, J. (2008), “Fuzzy Regression Analysis by Support Vector Learning Approach” **IEEE Transactions on Fuzzy Systems**, 16(2), pp. 428-441.
- [60] Hassanpour H., Maleki H. R. and Yaghoobi M. A. (2009), “A goal programming approach to fuzzy linear regression with non-fuzzy input and fuzzy output data” **Iranian Journal of Fuzzy Systems**, 6(2), pp. 1-6.
- [61] Hasanpour H., Maleki H. R. and Yaghoubi M. A. (2010), “Fuzzy linear regression model with crisp coefficients: a goal programming approach” **Iranian Journal of Fuzzy Systems**, 7(2), pp. 19-39.
- [62] Hassanpour H., Maleki H. R. and Yaghoobi M. A. (2011), “A goal programming approach to fuzzy linear regression with fuzzy input–output data” **Soft Computing**, 15(8), pp. 1569-1580.
- [63] Hesamian, G., Akbari, M. G., and Asadollahi, M. (2017), “ Fuzzy semi-parametric partially linear model with fuzzy inputs and fuzzy outputs” **Expert Systems with Applications**, 71, pp. 230-239.
- [64] Hesamian, G., Akbari, M. G. ,and Asadollahi M. (2017), “Fuzzy semi-parametric partially linear model with fuzzy inputs and fuzzy outputs“ **Expert Systems with Applications**, 71, pp. 230-239.
- [65] Hesamian G. and Akbari, M. G. (2019), “Fuzzy Lasso regression model with exact explanatory variables and fuzzy responses” **International Journal of Approximate Reasoning**, 115, pp. 290-300.
- [66] Hesamian G. and Akbari M. G. (2019), “Fuzzy Lasso regression model with exact explanatory variables and fuzzy responses” **International Journal of Approximate Reasoning**, 115, pp. 290-300.
- [67] Hoerl, A. E. and Kennard R. W. (1970), “Ridge regression: Biased estimation for nonorthogonal problems“ **Technometrics**, 12(1) pp. 55-67.

- [68] Hojati M., Bector C.R. and Smimou K. (2005), "A simple method for computation of fuzzy linear regression" **European Journal of Operational Research**, 166(1), pp. 172-184.
- [69] Hojati M. , Bector C. R. ,and Smimou, K. (2005), "A simple method for computation of fuzzy linear regression" **European Journal of Operational Research**, 166(1), pp. 172-184.
- [70] Hong, D.H. and Yi H.C. (2004), "A note on fuzzy regression model with fuzzy input and output data for manpower forecasting" **Fuzzy Sets and Systems**, 138(2), pp. 301-305.
- [71] Hose, D., and Hanss, M. (2018), "Fuzzy linear least squares for the identification of possibility regression models" **Fuzzy Sets and Systems**, 367, pp. 82-95.
- [72] Hosseinzadeh E., Hassanpour H. and Arefi M. (2016), "A weighted goal programming approach to estimate the linear regression model in full quasi type-2 fuzzy environment" **International Journal of Fuzzy Systems**, 30, pp. 1311-1317.
- [73] Huang, J. , Breheny, P. , Ma, S. ,and Zhang, C-H. (2010). "The Mnet Method for Variable Selection" (**Unpublished**) **Technical Report**, 26 .
- [74] James, W. and Stein, C. (1961), "Estimation with Quadratic Loss" **University of California Press, Berkeley, Calif.**, pp. 361-379.
- [75] Jung H.Y., Yoon J.H. and Choi S.H. (2015), "Fuzzy linear regression using rank transform method" **Fuzzy Sets and Systems**, 274, pp. 97-108.
- [76] Kannan, K. S. and Manjo, K. (2015), "Outlier detection in multivariate data" **Applied Mathematical Sciences**, 9(47), pp. 2317-2324.
- [77] Kao C. and Chyu C.L. (2002), "A fuzzy linear regression model with better explanatory power" **Fuzzy Sets and Systems**, 126(3), pp. 401-409.
- [78] Kao C. and Chyu C.L. (2003), "Least-squares estimates in fuzzy regression analysis" **European Journal of Operational Research**, 148(2), pp. 426-435.
- [79] Kashani M., Arashi M., Rabiei M.R. (2017), "Jackknife-after-Bootstrap in Fuzzy Regression Modeling" **17th conference on fuzzy systems and 15th conference on intelligent systems**.

-
- [80] Kelkinnama M. and Taheri S.M. (2012), "Fuzzy least-absolutes regression using shape preserving operations, **Information Sciences**, 214, pp. 105-120.
- [81] Kelkinnama, M. and Taheri, S. M. (2012), "Fuzzy least-absolutes regression using shape preserving operations" **Information Sciences**, 214, pp. 105-120.
- [82] Kim B. and Bishu, R.R. (1998), "Evaluation of fuzzy linear regression models by comparing membership functions" **Fuzzy Sets and Systems**, 100(1-3), pp. 343-352.
- [83] Kim I.K., Lee W., Yoon J.H. and Choi S.H. (2016), "Fuzzy regression model using trapezoidal fuzzy numbers for re-auction data" **International Journal of Fuzzy Logic and Intelligent Systems**, 16, pp. 72-80.
- [84] Lee H.T. and Chen S.H. (2001), "Fuzzy regression model with fuzzy input and output data for manpower forecasting" **Fuzzy Sets and Systems**, 119(2), pp. 205-213.
- [85] Lee W.J., Jung H.Y, Yoon J.H., Choi S.H. (2015), "The statistical inferences of fuzzy regression based on bootstrap techniques" **Soft Comp.**, 19, pp. 883-890.
- [86] Li J., Zeng W., Xie J. and Yin Q. (2016) "A new fuzzy regression model based on least absolute deviation" **Engineering Applications of Artificial Intelligence**, 52, pp. 54-64
- [87] Lin J. G., Zhuang Q. Y. and Huang, C. (2012), "Fuzzy statistical analysis of multiple regression with crisp and fuzzy covariates and applications in analyzing economic data of China" **Computational Economics**, 39(1), pp. 29-49.
- [88] Liu H-T., Wang J., He Y-L., Aamir, R. ,and Ashfaq R. (2017) "Extreme learning machine with fuzzy input and fuzzy output for fuzzy regression" **Neural Compute & Application**, 28(11), pp. 3465-3476.
- [89] Liu H. T., Wang J., He Y. L. and Ashfaq R. A. R. (2017). "Extreme learning machine with fuzzy input and fuzzy output for fuzzy regression" **Neural Computing and Applications**, 28(11), pp. 3465-3476
- [90] Lu J. and Wang R. (2009), "An enhanced fuzzy linear regression model with more flexible spreads" **Fuzzy Sets and Systems**, 160(17), pp. 2505-2523.

- [91] Mendel J.M. (2014), "On a novel way of processing data that uses fuzzy sets for later use in rule-based regression and pattern classification" **International Journal of Fuzzy Logic and Intelligent Systems**, 14, pp. 1-7.
- [92] Ming M., Friedman M. and Kandel A. (1997), "General fuzzy least squares" **Fuzzy Sets and Systems**, 88(1), pp. 107-118.
- [93] Mohammadi J. and Taheri S. (2004), "Pedomodels fitting with fuzzy least squares regression" **Iranian Journal of Fuzzy Systems**, 1(2), pp. 45-61.
- [94] Modarres M., Nasrabadi E. and Nasrabadi M. M. (2005), "Fuzzy linear regression models with least square errors" **Applied Mathematics and Computation**, 163(2), pp. 977-989.
- [95] Namdari M., Yoon J.H., Abadi A., Taheri S.M. and Choi S.H. (2015), "Fuzzy logistic regression with least absolute deviations estimators" **Soft Computing**, 19, pp. 909-917.
- [96] Nasrabadi M. M. and Nasrabadi E. (2004), "A mathematical-programming approach to fuzzy linear regression analysis" **Applied Mathematics and Computation**, 155(3), pp. 873-881.
- [97] Nasrabadi, M. M., Nasrabadi, E. ,and Nasrabady, A. R. (2005), "Fuzzy linear regression analysis: a multi-objective programming approach" **Applied Mathematics and Computation**, 163(1), pp. 245-251.
- [98] Nasrabadi M. M., Nasrabadi E., and Nasrabady A. R. (2005), "Fuzzy linear regression analysis: a multi-objective programming approach" **Applied Mathematics and Computation**, 163(1), pp. 245-251.
- [99] Nguyen H. T. (1978), "A note on the extension principle for fuzzy sets" **Journal of Mathematical Analysis and Applications**, 64(2), pp. 369-380.
- [100] Özelkan E.C. and Duckstein L. (2000), "Multi-objective fuzzy regression: a general framework" **Computers and Operations Research**, 27(7), pp. 635-652.
- [101] Parchami A. and Mashinch M. (2007), "Fuzzy estimation for process capability indices" **Information Science**, 177, pp. 1452-1462.

- [102] Parchami A. and Mashinch M. (2011), “Measuring process capability index with fuzzy data and compare it with other fuzzy process indices“ **Expert Systems with Applications**, 38, pp. 6452-6457.
- [103] Peters G. (1994), “Fuzzy linear regression with fuzzy intervals” **Fuzzy Sets and Systems**, 63(1), pp. 45-55.
- [104] Petit-Renaud, S. and Dencœux, T. (2004), “Nonparametric regression analysis of uncertain and imprecise data using belief functions“, **International Journal of Approximate Reasoning**, 35(1), pp. 1-28.
- [105] Pourahmad S., Ayatollahi M., Taheri S.M. and Habib Agahi Z. (2011), “Fuzzy logistic regression based on the least squares approach with application in clinical studies” **Computers and Mathematics with Applications**, 62, pp. 3353-3365.
- [106] Pourghasemi, H. R., Pradhan, B., and Gokceoglu, C. (2012), “Application of fuzzy logic and analytical hierarchy process (AHP) to landslide susceptibility mapping at Haraz watershed, Iran” **Natural hazards**, 63(2), pp. 965-996.
- [107] Rabiei M. R., Arashi M. and Farrokhi, M. (2019). Fuzzy ridge regression with fuzzy input and output. *Soft Computing*, 23(23), pp. 12189-12198.
- [108] Rabiei M.R., Arghami N.R., Taheri S.M. and Sadeghpour B. (2014), “Least squares approach to regression modeling in full interval-valued fuzzy environment” **Soft Computing**, 18, pp. 2043-2059.
- [109] Rabiei M.R., Taheri S.M. and Arghami N.R. (2015), “A linear-programming approach to interval-valued fuzzy regression analysis” **International Journal of Intelligent Technologies and Applied Statistics**, 8(3), pp. 171–203.
- [110] Rabiei M. R., Arashi M. and Farrokhi, M. (2019). Fuzzy ridge regression with fuzzy input and output. *Soft Computing*, 23(23), pp. .12189-12198
- [111] Rawlings, J. O. , Pantula, S. G. , and Dickey, D. A. (2001), “Applied regression analysis: a research tool” **Springer Science & Business Media**.
- [112] Redden D.T. and Woodall W.H. (1996), “Further examination of fuzzy linear regression” **Fuzzy Sets and Systems**, 79(2), pp. 203-211.
- [113] Redden D.T. and Woodall W.H. (1996) , “Further examination of fuzzy linear regression” **Fuzzy Sets and Systems**, 79(2) , pp. 203-211.

- [114] RodriguezFdez I., Mucientes M. and Bugarin A. (2016), "FRULER: Fuzzy rule learning through evolution for regression" **Information Sciences**, 354, pp. 1-18.
- [115] Sadeghpour Gildeh, B. and Gien, D. (2002), "A goodness of fit index to reliability analysis in fuzzy model" **In 3rd WSEAS international conference on fuzzy sets and fuzzy systems**, Iterlaken, Switzerland, pp. 11-14.
- [116] Sakawa M. and Yano H. (1992), "Multiobjective fuzzy linear regression analysis for fuzzy input-output data" **Fuzzy Sets and Systems**, 47(2), pp. 173-181.
- [117] Saleh A.M.E. (2006), **Theory of Preliminary Test and Stein-Type Estimation with Applications**, John Wiley, New York.
- [118] Saleh, A.K.M.E. (2006), "Theory of Preliminary Test and Stein-Type Estimation with Applications" **Wiley, New Jersey**, .
- [119] Saleh, A.K.M.E. and Arashi, M. and Kibria, B.M.G. (2019), "Theory of Ridge Regression Estimation with Applications" **Wiley Series in Probability and Statistics, Wiley, USA**, .
- [120] Savic D.A. and Pedrycz W. (1991), "Evaluation of fuzzy linear regression models" **Fuzzy Sets and Systems**, 39(1), pp. 51-63.
- [121] Sharma, L. K., Vishal, V., and Singh, T. N. (2017), "Developing novel models using neural networks and fuzzy systems for the prediction of strength of rocks from key geomechanical properties" **Measurement**, 102, pp. 158-169.
- [122] Sohn, B.-Y. (2005), "Robust fuzzy linear regression based on M-estimators" **J. Appl. Math. & Computing**, 18, pp. 591-601.
- [123] Stine, C. M. (1981), "Estimation of the Mean of a Multivariate Normal Distribution" **The Annals of Statistics**, 9(6), pp. 1135-1151.
- [124] Stein C.M. (1981), "Estimation of the mean of a multivariate normal distribution" **The annals of Statistics**, pp. 1135-1151.
- [125] Su Z., Wang Y. and Wang P. (2013), "Parametric regression analysis of imprecise and uncertain data in the fuzzy belief function framework" **International Journal of Approximate Reasoning**, 54, pp. 1217-1242.

-
- [126] Taheri S.M., Salmani F., Abadi A. and Alavi H. (2016), "A transition model for fuzzy correlated longitudinal responses" **Journal Intelligent and Fuzzy Systems**, 30, pp. 1265-1273.
 - [127] Tanaka, Hideo and Uejima, S. and Asai, K. (1982), "Linear regression analysis with fuzzy model" **IEEE Transactions on Systems, Management, and Cybernetics**, 12, pp. 903-907.
 - [128] Tanaka H. (1987), "Fuzzy data analysis by possibilistic linear models" **Fuzzy Sets and Systems**, 24(3), pp. 363-375.
 - [129] Tanaka H., Hayashi I. and Watada J. (1989), "Possibilistic linear regression analysis for fuzzy data" **European Journal of Operational Research**, 40(3), pp. 389-396.
 - [130] Tanaka H. and Ishibuchi H. (1991), "Identification of possibilistic linear systems by quadratic membership functions of fuzzy parameters" **Fuzzy Sets and Systems**, 41(2), pp. 145-160.
 - [131] Tanaka H., Ishibuchi H. and Yoshikawa S. (1995), "Exponential possibility regression analysis" **Fuzzy Sets and Systems**, 69(3), pp. 305-318.
 - [132] Tanaka H. and Lee H. (1999), "Exponential possibility regression analysis by identification method of possibilistic coefficients" **Fuzzy Sets and Systems**, 106(2), pp. 155-165.
 - [133] Taylor, K. E. (2001), "Summarizing multiple aspects of model performance in a single diagram" **Journal of Geophysical Research: Atmospheres**, 106(D7), pp. 7183-7192.
 - [134] Tibshirani R. (1996), "Regression shrinkage and selection via the lasso" **Journal of the Royal Statistical Society: Series B (Methodological)**, 58(1), pp. 267-288.
 - [135] Torabi H., Behboodian J. (2007), "Fuzzy least-absolutes estimates in linear models" **Communications in Statistics—Theory and Methods**, 36(10), pp. 1935-1944.
 - [136] Tsujinishi, D., and Abe, S. (2003), "Fuzzy least squares support vector machines for multiclass problems", **Neural Networks**, 16(5-6), pp. 785-792.
 - [137] Wang H.F. and Tsaur R.C. (2000), "Resolution of fuzzy regression model" **European Journal of Operational Research**, 126(3), pp. 637-650.

- [138] Wang, N., Zhang, W. X., and Mei, C. L. (2007), “Fuzzy nonparametric regression based on local linear smoothing technique” **Information sciences**, 177(18), pp. 3882-3900.
- [139] Wu, H. C. (2003), “Linear regression analysis for fuzzy input and output data using the extension principle“ **Computers & Mathematics with Applications**, 45(12), pp. 1849-1859.
- [140] Wu H.C. (2003), “Fuzzy estimates of regression parameters in linear regression models for imprecise input and output data” **Computational Statistics & Data Analysis**, 42(1-2), pp. 203-217.
- [141] Wu, H. C. (2011), “The construction of fuzzy least squares estimators in fuzzy linear regression models“ **Expert Systems with Applications**, 38(11), pp. 13632-13640.
- [142] Xu R., Li C. (2001), “Multidimensional least-squares fitting with a fuzzy model” **Fuzzy Sets and Systems**, 119(2), pp. 215-223.
- [143] 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(1), pp. 49-60.
- [144] Yang M.S. and Lin T.S. (2002), “Fuzzy least-squares linear regression analysis for fuzzy input-output data” **Fuzzy Sets and Systems**, 126(3), pp. 389-399.
- [145] Yang M.S. and Liu H.H. (2003), “Fuzzy least-squares algorithms for iterative fuzzy linear regression models” **Fuzzy Sets and Systems**, 135, pp. 305-316.
- [146] Yüzbaşı B., Arashi M. and Ahmed S. E. (2019), “Shrinkage estimation strategies in generalized ridge regression models under low/high-dimension regime” **International Statistical Review**, 96(456), pp. 1348–1360.
- [147] Yüzbaşı, Bahadır and Arashi, M. and Ahmed, S. E. (2017), “Shrinkage Estimation Strategies in Generalized Ridge Regression Models Under Low/High-Dimension Regime“ **International Statistical Review**, 96(456), pp. 1348-1360.
- [148] Yüzbaşı, B. and Arashi, M. (2019), “Double shrunken selection operator“ **Communications in Statistics-Simulation and Computation**, 48(3), pp. 666-674.
- [149] Zadeh L. A. (1965), “Fuzzy sets” **Inform and Control**, 8, pp. 338-353.

- [150] Zeng W., Feng Q. and Li J. (2017), “Fuzzy least absolute linear regression” **Applied Soft Computing**, 52, pp. 1009-1019.
- [151] Zhou, J. , Zhang, H. , Gu, Y. ,and Pantelous, A. A. (2018). “Affordable levels of house prices using fuzzy linear regression analysis: the case of Shanghai” **Soft Computing**, 22(16), pp. 5407-5418.
- [152] Zhao, H. , Zheng, J. , Dang, W. ,and Song, Y. (2020), “Semi-supervised broad learning system based on manifold regularization and broad network” **IEEE Transactions on Circuits and Systems I: Regular Papers**, 67(3), pp. 983-994.
- [153] Zou, H. and Hastie, T. (2005), “Regularization and variable selection via the elastic net” **Journal of the royal statistical society: series B (statistical methodology)**, 67(2), pp. 301-320.
- [154] Zou, H. and Hastie, T. (2005), “Regularization and variable selection via the elastic net” **Journal of the royal statistical society: series B (statistical methodology)**, 67(2), pp. 301-320.

Aabstract

In the existing views of fuzzy regression analysis, the improvement of the model parameter estimators has been only from the perspective of correction of the existing meters and improving the problems of the distances and previous meters. While in classical (non-fuzzy) regression, the estimators are improved in order to achieve the optimality of the comparison criteria and the goodness of fit is also done by manipulating and generalizing the estimators themselves.

Therefore, in this study, we aim to improve the fuzzy optimization criteria by improving the estimators in fuzzy regression models, and not the desired meter. In this regard, we propose the use of shrinkage estimators, resampling methods, and penalized methods in fuzzy regression models and show how fuzzy optimization criteria can be improved by using fuzzy improved estimators (shrinkage and penalized methods) or retrieval of fuzzy data.

Key words: Fuzzy regression; Fuzzy Stein shrinkage estimator; Jackknife-after-Bootstrap; Penalized estimator.



Faculty Of Mathematical Sciences

PhD Thesis in Fuzzy statistical inference

Improved estimators in fuzzy regression models

By: Mojtaba Kashani

Supervisors:

Dr. Mohammad Arashi
Dr. Mohammad Reza Rabiei

April 2021