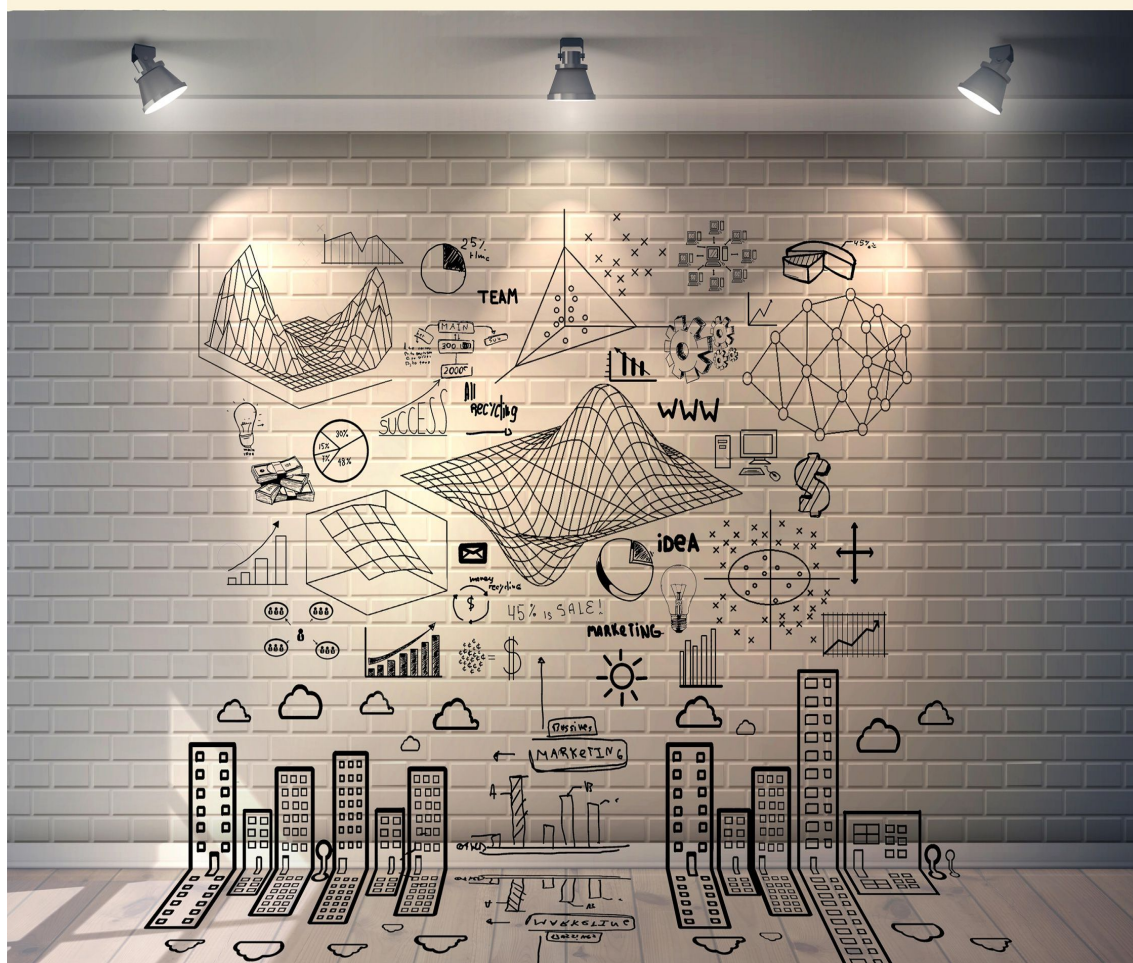


مجموعه مقالات

سومین سمینار آمار فضایی و کاربردهای آن



به نام خدا



سومین سمینار آمار فضایی و کاربردهای آن

مجموعه مقالات

۱ الی ۲ آبان ۱۳۹۸

دانشگاه زنجان

مجموعه مقالات سومین سمینار آمار فضایی و کاربردهای آن

چاپ: اول آبان ماه ۱۳۹۸
آدرس دبیرخانه: زنجان، دانشگاه زنجان، دانشکده علوم، گروه آمار
تلفن: ۰۲۴ - ۳۳۰۵۴۰۴۲
پست الکترونیک: spatial3@znu.ac.ir
آدرس سایت: spatial3.znu.ac.ir

پیش‌گفتار

سپاس بی‌کران خدای منان را که هستی‌مان بخشید، به طریق علم و دانش رهنمونمان شد، به همنشینی رهروان علم و دانش مفتخرمان نمود و خوشه‌چینی از علم و دانش را روزیمان ساخت. ضمن عرض خیر مقدم به همه شرکت‌کنندگان گرامی و تشکر از حضور آنان، بسیار خرسندیم که در راستای خدمت به جامعه علمی آمار کشور، اعتلای علمی این مرز و بوم و به مناسبت گرامی‌داشت روز ملی آمار و برنامه‌ریزی، سومین سمینار آمار فضایی و کاربردهای آن را با همکاری انجمن آمار ایران و قطب علمی تحلیل داده‌های وابسته فضایی-زمانی دانشگاه تربیت مدرس و به میزبانی گروه آمار دانشگاه زنجان برگزار نماییم. امید است برنامه‌های سمینار که حاصل تلاش‌های پیوسته و موثر کمیته‌های علمی و اجرایی است، مورد رضایت واقع شود. تمام تلاش کمیته برگزارکننده بر آن بوده که فرصتی ایجاد شود تا همه پژوهشگران، متخصصان و علاقه‌مندان به این شاخه از علم آمار به ارائه آخرین یافته‌های پژوهشی خود پرداخته و زمینه توسعه و استفاده جدی‌تر این شاخه از آمار کاربردی را در سایر حوزه‌های علمی و اجتماعی فراهم آورند. امید است توانسته باشیم با لطف و عنایت پروردگار و همراهی و همدلی شرکت‌کنندگان گرامی، این سمینار را به نحو مطلوب برگزار کرده باشیم. با این وجود اگر کمبودهایی در برگزاری سمینار بوده است آن را بر ما ببخشایید و با اعلام آنها به کمیته برگزاری، امکان برگزاری بهتر سمینارهای آتی آمار فضایی و کاربردهای آن را فراهم آورید. در پایان از همه کسانی که ما را در برگزاری سمینار یاری نموده‌اند کمال سپاس و امتنان را داریم. بخصوص از داوران گرامی، قطب علمی تحلیل داده‌های وابسته فضایی-زمانی دانشگاه تربیت مدرس و همکاران هیات علمی، دانشجویان گروه آمار و معاونت محترم پژوهشی، ریاست محترم دانشگاه و کارکنان مختلف دانشگاه زنجان که برگزاری این سمینار بدون مساعدت‌های بی‌دریغ و همه‌جانبه آن‌ها امکان‌پذیر نبود، قدردانی و تشکر می‌نماییم. امید است با برگزاری این سمینار گامی هر چند کوچک در راستای ارتقای دانش و تجربه در حوزه آمار فضایی کشور برداشته باشیم.

دبیر سومین سمینار آمار فضایی و کاربردهای آن

دانشگاه زنجان

۱ الی ۲ آبان ۱۳۹۸

فهرست همکاران سمینار:



انجمن آمار ایران



قطب علمی تحلیل داده‌های وابسته فضایی-زمانی

فهرست حامیان سمینار:



پایگاه استنادی علوم جهان اسلام



مرکز منطقه‌ای اطلاع‌رسانی علوم و فناوری

اعضای کمیته علمی

۱. دکتر علی محمدیان مصمم (دبیر کمیته علمی) دانشگاه زنجان
۲. دکتر علی آقامحمدی دانشگاه زنجان
۳. دکتر حسین باغیشنی دانشگاه صنعتی شاهرود
۴. دکتر امید کریمی دانشگاه سمنان
۵. دکتر افشین فلاح دانشگاه بین المللی امام خمینی
۶. دکتر عباس رسولی دانشگاه زنجان
۷. دکتر محسن محمدزاده دانشگاه تربیت مدرس

اعضای کمیته اجرایی

۱. دکتر علی آقامحمدی (دبیر کمیته اجرایی) دانشگاه زنجان
۲. دکتر فرض اله میرزاپور دانشگاه زنجان
۳. دکتر عباس رسولی دانشگاه زنجان
۴. دکتر علی محمدیان مصمم دانشگاه زنجان
۵. دکتر هادی امامی دانشگاه زنجان
۶. دکتر امید خادم‌نوع دانشگاه زنجان
۷. دکتر مهری جوانیان دانشگاه زنجان

فهرست مقالات

عنوان	صفحه
نسرين ابراهيمي و محسن محمدزاده	
مدل‌بندی داده‌های بقای فضایی با تابع مفصل زوجی کلايتون	۱
اسحاق الماسی، مهدی امیدی	
پیشگویی فضایی میدان تصادفی ناگوسی با استفاده از قضیه تصویر	۷
جمیل اونق، حسین باغیشنی	
رگرسیون توزیعی فضایی با دم‌های نیمه‌سنگین	۱۱
سپیده تشرقی، امید کریمی، فاطمه حسینی	
تحلیل مدل‌های سری زمانی شمارشی فضایی	۲۲
عباس توسلی، پدا... واقعی و علیرضا ناظمی	
برآورد تغییرنگار با استفاده از روش رگرسیون بردار پشتیبان	۳۲
فاطمه حسینی، امید کریمی	
تحلیل بیزی تقریبی مدل‌های رگرسیونی فضایی	۴۰
سحر خلفی، علی محمدیان مصمم، علی آقامحمدی	
مقدمه‌ای بر تقریب لاپلاس آشیانه‌ای جمع بسته: رهیافت سریع و دقیق بیزی	۵۰
زهرا رحیمیان آزاد، افشین فلاح، زهرا زارع‌پور یکنمی	
میانگین‌گیری بیزی مدل‌های رگرسیونی فضایی تحت توزیع گاوسی وارون	۵۶
سمیرا زحمتکش، محسن محمدزاده	
مدل‌بندی توأم بیزی داده‌های فضایی-زمانی با گمشدگی غیرقابل چشم‌پوشی	۶۵
فاطمه زنجیران، کیومرث مترجم	
اثر الگوی جغرافیایی بر رضایت مشتری	۷۴
سمیرا سعادت‌ی و محسن محمدزاده	
تحلیل داده‌های فضایی-زمانی بر اساس تجزیه تاکر تانسور کوواریانس	۸۵
علی شهنواز، علی محمدیان مصمم، محمد حسن بهزادی	
مقایسه کارایی آماری و محاسباتی روش‌های برآوردیابی در مدل‌های نیمه‌طیفی فضایی-زمانی	۹۳
لیلا صالحی و محسن محمدزاده	

پنج مجموعه مقالات

مقایسه دو تابع درستمایی مرکب در برآورد مدل های آمیخته خطی تعمیم یافته فضایی ۱۰۹
مسعود عظیمیان، محمد مرادی

مروری بر نمونه گیری مجموعه رتبه دار در جوامع فضایی ۱۱۹
فتانه فخری، حسین باغیشنی، بهزاد تخم چی

پیشگویی فضایی سه بعدی به منظور برآورد مقاومت فشاری تک محوری در یکی از مخازن نفتی ایران ۱۲۴
زهره فضلعلی و کیومرث مترجم

بررسی اثر موقعیت مکانی هتل ها بر میزان رضایت مندی مشتریان ۱۳۳
زهره زارع پور یکنمی، افشین فلاح

تحلیل کواریانس گاوسی وارون برای داده های فضایی ۱۴۲
کبری قلی زاده، محسن محمدزاده

مدل بندی بیزی داده های نرخ فضایی با مدل رگرسیون بتا فضایی ۱۴۹
الهه لطفیان و محسن محمدزاده

ارتقای الگوریتم چند هدفه ازدحام ذرات برای نمونه گیری بهینه فضایی ۱۵۸
یدالله محرابی، امیر کاوسی، بهزاد مهکی، پریسا پورنظری

مدل های توام نقشه بندی بیماری ها ۱۷۰
علی محمدیان مصمم، پریا عیوضی

مدل سازی داده های فضایی نایستا ۱۸۲
علی محمدیان مصمم، زهره مقیمی

تحلیل داده های فضایی- زمانی در حضور مقادیر گمشده ۱۸۶
علی محمدیان مصمم، الناز عباسی

تحلیل بیزی داده های گسسته فضایی- زمانی ۱۹۰
سرور موسوی، یونس خسروی، عباسعلی زمان، مهسا دادشی

پایش کیفیت آب های سطحی با استفاده از تکنیک های آمار فضایی: مطالعه موردی استان زنجان ۱۹۶
حسین والی، امید کریمی، فاطمه حسینی

تحلیل زمانی و فضایی- زمانی داده های ریسک سلامت در تصادفات جاده ای استان قم ۲۰۰
حسین ویسی پور و محمد مرادی

مروری بر طرح های نمونه گیری فضایی ساخته شده از دنباله هالتون ۲۱۴
فاطمه یزدانی، مجید عظیم محسنی، مهناز خلفی

سومین سمینار آمار فضایی و کاربردهای آن شمش

استفاده از رتبه‌ها در تحلیل ممیزی فضایی داده‌های چند متغیره ۲۲۱



مدل‌بندی داده‌های بقای فضایی با تابع مفصل زوجی کلایتون

نسرین ابراهیمی^۱، محسن محمدزاده

گروه آمار، دانشگاه تربیت مدرس

چکیده: بررسی تأثیر عوامل مخاطره مختلف بر زمان بقای آزمودنی‌ها همواره مورد توجه محققان در حوزه تحلیل داده‌های بقا بوده است. هنگامی که بین زمان‌های بقا وابستگی وجود داشته باشد، این وابستگی به عنوان یک عامل مخاطره پنهان عمل می‌کند و بررسی تأثیر آن بر زمان‌های بقا از طریق مدل‌های رگرسیونی معمول بقا نظیر مدل خطرهای متناسب کاکس امکان‌پذیر نیست. یکی از روش‌های مدل‌بندی زمان‌های بقا با لحاظ کردن وابستگی بین آن‌ها، استفاده از توابع مفصل است. در این مقاله تابع بقا با استفاده از توابع بقای توأم دو به دو از طریق تابع مفصل زوجی کلایتون مدل می‌شود. تابع درستنمایی مرکب برای برآورد پارامترهای وابستگی به کار می‌رود و عملکرد مدل معرفی شده از طریق مطالعه شبیه‌سازی مورد بررسی قرار می‌گیرد.

واژه‌های کلیدی: داده‌های بقای فضایی، تابع بقای توأم، تابع درستنمایی مرکب، مدل‌های حاشیه‌ای، تابع مفصل کلایتون. **کد موضوعی:** بند ریاضی (۲۰۱۰): 62N01.

۱ مقدمه

هدف عمده بسیاری از مطالعات صورت گرفته در حوزه تحلیل داده‌های بقا، تعیین تأثیر عوامل مخاطره مختلف روی زمان بقا است. در بسیاری از موارد، عوامل خطر ناشناخته یا غیر قابل اندازه‌گیری نیز روی زمان بقا تأثیرگذار هستند و موجب وابستگی بین زمان‌های بقا می‌شوند که نادیده گرفتن این وابستگی در مدل، نتایج گمراه‌کننده‌ای را به همراه خواهد داشت. یکی از این عوامل ناشناخته تأثیرگذار بر داده‌های بقا، موقعیت فضایی مشاهده مورد نظر است، به‌طوری که هر داده بقا برحسب موقعیت خود در فضای مورد مطالعه، با داده‌های دیگر وابستگی فضایی دارد. اما این اثر فضایی به‌طور مستقیم قابل مشاهده یا شناسایی نیست. تحلیل و بررسی داده‌های بقای فضایی همواره مورد توجه محققین علوم زیستی در مطالعات همه‌گیرشناسی و آزمایش‌های بالینی خانوادگی بوده است. تحلیل داده‌های بقای فضایی مستلزم مدل‌بندی تابع بقای توأم این داده‌هاست. در حالت کلی دو رهیافت عمده برای مدل‌بندی داده‌های بقای فضایی با لحاظ کردن وابستگی آن‌ها، استفاده از

^۱ نام و ایمیل ارائه دهنده مقاله: نسرین ابراهیمی، nasrin.ebrahimi@modares.ac.ir

مدل‌های شکنندگی و توابع مفصل است. مدل‌های شکنندگی، وابستگی فضایی بین آزمودنی‌ها را در قالب یک اثر تصادفی پنهان به نام شکنندگی در مدل لحاظ می‌کنند. برآورد پارامترهای حاصل از این مدل، دارای تفسیرهایی به ازای یک مقدار خاص شکنندگی هستند و در مواردی که استنباط‌های درون‌خوشه‌ای مدنظر باشند مورد استفاده قرار می‌گیرند.

لی و ریان (۲۰۰۲) مدل بقای فضایی را با اضافه کردن اثرات تصادفی فضایی به مدل نیم‌پارامتری خطرهای متناسب معرفی کردند و این مدل را در برازش به داده‌های مربوط به آسم کودکان به کار بردند. **لی و ریان (۲۰۰۲)** یک میدان تصادفی گاوسی را به عنوان توزیع اثرات تصادفی فضایی در نظر گرفتند. **مترجم و همکاران (۱۳۹۴)** این مدل را برای میدان تصادفی ناگاوسی مورد مطالعه قرار دادند و شناسایی‌پذیری پارامترهای مدل را در این حالت اثبات کردند. **بانرجی و دی (۲۰۰۵)** از تکنیک مدل‌بندی شکنندگی در قالب مدل نیم‌پارامتری بخت‌های متناسب استفاده کردند. **پن و همکاران (۲۰۱۴)** مدل نیم‌پارامتری خطرهای متناسب را با شکنندگی‌های CAR برای داده‌های سانسوریده بازه‌ای به کار بردند. یکی از ابزارهای قوی برای لحاظ کردن وابستگی داده‌ها در مدل استفاده از توابع مفصل است. این توابع مدلی را برای بررسی ساختار ارتباط بین متغیرها ارائه می‌دهند که بر اساس آن تمام خصوصیات وابستگی میان متغیرها قابل تبیین است (**امیدی و محمدزاده، ۱۳۹۰**). با توجه به این ویژگی، می‌توان توابع مفصل را برای مدل کردن تابع بقای توأم بر اساس توابع بقای حاشیه‌ای مفروض به کار برد. **کلایتون (۱۹۷۸)** اولین بار یک مدل پیوند دومتغیره را در تحلیل داده‌های بقای وابسته به کار برد. در مدل پیشنهادی او یک تابع مفصل برای مدل‌بندی زمان‌های بقای دومتغیره وابسته به کار رفته است. **امیدی و محمدزاده (۱۳۹۰)** و **شیائو (۲۰۰۶)** از توابع مفصل برای مدل‌بندی داده‌های خشکسالی استفاده کردند. **لی و لین (۲۰۰۶)** با فرض آن‌که زمان‌های بقا به‌طور حاشیه‌ای از مدل نیم‌پارامتری خطرهای متناسب پیروی می‌کنند، از تابع مفصل گاوسی برای مدل‌بندی ساختار همبستگی فضایی داده‌های بقای زمین‌آماري با یک ساختار همبستگی خاص استفاده کردند. **ژو و همکاران (۲۰۱۵)** از رهیافت بیز ناپارامتری برای مدل‌بندی تابع خطر داده‌های بقای زمین‌آمار استفاده کردند و تابع مفصل گاوسی را برای مدل‌بندی ساختار همبستگی فضایی آن‌ها به کار بردند. همچنین نسخه بیزی مدل معرفی شده توسط **لی و لین (۲۰۰۶)** را با در نظر گرفتن پیشین تکه‌ای‌نمایی برای تابع خطر پایه، ارائه کردند. **پرن و همکاران (۲۰۱۷)** روش‌شناسی توابع مفصل ارشمیدسی را برای مدل‌بندی داده‌های بقای خوشه‌بندی شده بسط دادند.

در این مقاله با در نظر گرفتن مدل حاشیه‌ای خطرهای متناسب برای زمان‌های بقای وابسته فضایی، از تابع مفصل کلایتون برای مدل‌بندی ساختار همبستگی آن‌ها استفاده می‌شود و یک روش دو مرحله‌ای برای برآورد پارامترهای مدل به کار می‌رود که در آن ابتدا پارامترهای رگرسیونی مدل به روش ماکسیمم درستنمایی جزئی برآورد می‌شوند و سپس در برآورد پارامتر ساختار همبستگی فضایی داده‌ها با استفاده از تابع درستنمایی مرکب به کار می‌روند.

۲ معرفی مدل

فرض کنید به ازای $i = 1, \dots, n$ متغیرهای T_i و C_i به ترتیب زمان‌های شکست و سانسور را برای آزمودنی i ام نشان دهند. در عمل T_i و C_i به‌طور بالقوه مشاهده نشده‌اند و داده‌های مشاهده شده عبارتند از: $X_i = \min(T_i, C_i)$ و $\Delta_i = I(T_i \leq C_i)$. فرض کنید $Z_i(t)$ بردار متغیرهای کمکی با بعد p باشد. همچنین مسیر^۱ متغیرهای کمکی تا زمان t با $\bar{Z}_i(t) = \{Z_i(s) | 0 \leq s \leq t\}$ نمایش داده شود و فرض کنید T_i به شرط $\bar{Z}_i(T_i)$ مستقل از C_i است. همچنین به شرط معلوم بودن $\bar{Z}_i(t)$ ، تابع نرخ خطر به صورت

$$\lambda(t|\bar{Z}_i(t)) = \lambda_0(t) \exp(\beta' Z_i(t)) \quad (1.2)$$

از مدل خطرهای متناسب پیروی می‌کند. معادله (۱.۲) یک مدل حاشیه‌ای برای هر T_i است و از آن‌جا که ضریب رگرسیونی β برای تمام i ها ثابت در نظر گرفته شده است، دارای تفسیری در سطح جامعه است. برای مدل‌بندی ساختار وابستگی T_i ها، فرض کنید که تابع توزیع توأم $F(t_i, t_j)$ با استفاده از توابع توزیع حاشیه‌ای $F(t_i)$ و $F(t_j)$ از طریق تابع مفصل کلایتون به صورت

$$F(t_i, t_j) = C(F(t_i), F(t_j)) = (F(t_i)^{-\alpha_{ij}} + F(t_j)^{-\alpha_{ij}} - 1)^{-1/\alpha_{ij}} \quad (۲.۲)$$

حاصل شود، که در آن $\alpha_{ij} \in (0, \infty)$ پارامتر تابع مفصل کلایتون است که میزان وابستگی دو متغیر T_j و T_i را اندازه می‌گیرد. نلسن (۲۰۰۵) ارتباط بین تابع مفصل بقا و تابع مفصل را در حالت کلی به صورت

$$S(t_i, t_j) = S(t_i) + S(t_j) - 1 + C(F(t_i), F(t_j)) \quad (۳.۲)$$

بیان می‌کند. بنابراین برای تابع مفصل کلایتون، تابع مفصل بقا یا به عبارت دیگر همان تابع بقای توأم به صورت

$$S(t_i, t_j) = S(t_i) + S(t_j) - 1 + (F(t_i)^{-\alpha_{ij}} + F(t_j)^{-\alpha_{ij}} - 1)^{-1/\alpha_{ij}} \quad (۴.۲)$$

است، که در آن $S(t_i)$ و $S(t_j)$ به ترتیب توابع بقای حاشیه‌ای T_i و T_j هستند که براساس مدل حاشیه‌ای خطرهای متناسب برآورد می‌شوند. بین پارامتر تابع مفصل کلایتون و ضریب همبستگی τ - کندال رابطه یک به یک به صورت $\alpha = 2\tau/(1 - \tau)$ وجود دارد (نلسن، ۲۰۰۵). در این مقاله ضریب همبستگی τ - کندال بین واحدهای i ام و j ام، تابعی نمایی از فاصله اقلیدسی بین دو واحد به صورت

$$\tau_{ij} \equiv \varphi(d_{ij}; \psi) = \exp(-\frac{d_{ij}}{\psi}) \quad (۵.۲)$$

در نظر گرفته شده است، که در آن d_{ij} فاصله اقلیدسی بین موقعیت‌های فضایی واحدهای i ام و j ام و ψ پارامتر وابستگی فضایی مدل است. برای برآورد مدل باید ضرایب رگرسیونی مدل حاشیه‌ای، β ، و پارامتر وابستگی فضایی، ψ ، برآورد شوند.

۳ برآورد مدل

برای برآورد پارامترهای مدل، از یک روش دو مرحله‌ای استفاده می‌شود. ابتدا ضرایب رگرسیونی مدل حاشیه‌ای خطرهای متناسب با روش ماکسیمم درستنمایی جزئی برآورد می‌شوند. سپس برآوردهای حاصل از این مرحله در تابع درستنمایی مرکب قرار می‌گیرند و در برآورد پارامتر وابستگی فضایی ψ به کار می‌روند.

مرحله اول: برآورد ضرایب رگرسیونی

در مدل رگرسیونی خطرهای متناسب کاکس هنگام برآورد پارامتر β تابع خطر پایه به‌عنوان یک پارامتر مزاحم نامتناهی-بعد ظاهر می‌شود، بنابراین روش‌های درستنمایی معمول نمی‌توانند در برآورد پارامتر رگرسیونی به کار روند. کاکس (۱۹۷۵) استنباط مبتنی بر تابع درستنمایی جزئی را در برآورد مدل‌های شامل پارامتر مزاحم نامتناهی-بعد معرفی کرد. برآورد ماکسیمم درستنمایی جزئی پارامتر β با ماکسیمم کردن تابع درستنمایی جزئی

$$L_p = \prod_{i=1}^n \left[\frac{\exp(\beta' Z_i(t))}{\sum_{j \in \mathcal{R}(X_i)} \exp(\beta' Z_j(t))} \right]^{\Delta_i} \quad (۱.۳)$$

نسبت به β حاصل می‌شود، که در آن $\mathcal{R}(X_i)$ مجموعه واحدهای در مخاطره تا زمان t است.

مرحله دوم: برآورد پارامتر وابستگی فضایی

برای برآورد پارامتر وابستگی مدل، از تابع درستنمایی مرکب استفاده شده است. تابع درستنمایی مرکب برای داده‌های بقا در حالت کلی به صورت

$$LC = \prod_{i \leq j} S(t_i, t_j)^{(1-\delta_i)(1-\delta_j)} S^{(1)}(t_i, t_j)^{\delta_i(1-\delta_j)} S^{(2)}(t_i, t_j)^{(1-\delta_i)\delta_j} f(t_i, t_j)^{\delta_i\delta_j}. \quad (2.3)$$

است، که در آن $S(t_i, t_j) = \frac{\partial^2}{\partial t_i \partial t_j} S(t_i, t_j)$ و $S^{(2)}(t_i, t_j) = \frac{\partial}{\partial t_j} S(t_i, t_j)$ ، $S^{(1)}(t_i, t_j) = \frac{\partial}{\partial t_i} S(t_i, t_j)$ است. برای مدل معرفی شده، با قرار دادن برآورد پارامتر رگرسیونی حاصل از مرحله قبل در تابع درستنمایی مرکب و ماکسیم کردن آن نسبت به ψ ، برآورد ماکسیم درستنمایی پارامتر وابستگی حاصل می‌شود.

۴ مطالعه شبیه‌سازی

برای شبیه‌سازی داده‌های بقای فضایی، ابتدا با فرض معلوم بودن پارامتر وابستگی فضایی، ۱۰۰ موقعیت فضایی از توزیع یکنواخت تولید و فاصله اقلیدسی این نقاط در قالب ماتریس فاصله محاسبه شد و با در نظر گرفتن مدل کواریوگرام نمایی به صورت

$$C(d) = \exp\left(-\frac{d}{\psi}\right) \quad (1.4)$$

ماتریس کواریانس مشاهدات به دست آمد. برای تولید داده بقا از مدل خطرهای متناسب، با استفاده از رابطه

$$S(t|Z) = \exp\left(-\Lambda_0(t) \exp(\beta'Z)\right) \quad (2.4)$$

در مدل خطرهای متناسب و با توجه به قضیه تبدیل انتگرال احتمال، می‌توان نتیجه گرفت که متغیر تصادفی

$$T = \Lambda_0^{-1} \left(-\frac{\ln S}{\exp(\beta'Z)} \right) \quad (3.4)$$

دارای تابع بقای (۲.۴) خواهد بود. حال با تولید ۱۰۰ مشاهده از توزیع $U(0, 1)$ به عنوان زمان بقای S و در نظر گرفتن هر تابع خطر پایه دلخواهی می‌توان مشاهداتی از مدل خطرهای متناسب شبیه‌سازی کرد. در این مقاله زمان‌های بقا با فرض تابع خطر پایه وایبل با پارامترهای $\lambda = 0.1$ و $\rho = 1$ با استفاده از رابطه

$$t|z = \left(\frac{-\ln S}{\lambda \exp(\beta'Z)} \right)^{\frac{1}{\rho}}$$

شبیه‌سازی شدند. با ضرب تجزیه چولسکی ماتریس کواریانس محاسبه شده در رابطه (۱.۴) در داده‌های بقای مستقل شبیه‌سازی شده و سپس حذف روند، داده‌های بقای فضایی با ساختار کواریانس موردنظر حاصل می‌شوند. متغیر تبیینی Z از توزیع برنولی شبیه‌سازی شد و نرخ سانسوری در حدود r درصد برای داده‌ها در نظر گرفته شد. این روند شبیه‌سازی ۱۰۰ بار تکرار شد. در هر بار تکرار برآورد ماکسیم درستنمایی پارامتر β محاسبه شد. نتایج این شبیه‌سازی در جدول ۱ مشاهده می‌شود. همان‌طور که ملاحظه می‌شود، برآوردهای حاصل از این روش برای پارامتر رگرسیونی β از دقت مناسبی برخوردار هستند. همان‌طور که انتظار می‌رود، با افزایش نرخ سانسور، دقت برآوردها تا حدودی کاهش می‌یابد. نکته قابل

جدول ۱: مقادیر واقعی پارامتر β به همراه برآورد و MSE آن

$r = 0.5$		$r = 0.2$		$r = 0.05$		مقدار واقعی β	مقدار معلوم ψ
MSE	برآورد β	MSE	برآورد β	MSE	برآورد β		
۰/۱۰۲۸	۰/۵۲۷۲	۰/۰۵۷۷	۰/۵۲۵۴	۰/۰۴۵۲	۰/۴۸۸۵	۰/۵	۰/۱
۰/۰۹۱	۰/۴۹۸۲	۰/۰۴۸۶	۰/۵۳۲۶	۰/۰۳۷۵	۰/۴۹۸۳	۰/۵	۰/۵
۰/۱۲۵	۰/۵۵۱۶	۰/۰۵۵۸	۰/۴۷۷۶	۰/۰۴۷۳	۰/۴۸۲۷	۰/۵	۲
۰/۱۱۸۶	۱/۵۱۶۱	۰/۰۹۴۳	۱/۵۰۶۶	۰/۰۸۴۹	۱/۵۷۸۳	۱/۵	۰/۱
۰/۱۰۲	۱/۵۴۰۱	۰/۰۸۹۸	۱/۵۱۳۳	۰/۰۵۵۴	۱/۵۲۱۶	۱/۵	۰/۵
۰/۰۸۹	۱/۴۴۰۵	۰/۱۰۱	۱/۴۰۹	۰/۰۹۹۶	۱/۳۸۱۱	۱/۵	۲
۰/۳۷۳۹	۳/۱۲۴۹	۰/۱۹۸۲	۳/۰۷۸۷	۰/۱۵۰۹	۳/۱۱۸۶	۳	۰/۱
۰/۳۶۰۸	۲/۹۲۶۲	۰/۲۴۹۶	۲/۹۸۴۴	۰/۲۹۱	۳/۰۱۶۲	۳	۰/۵
۰/۸۷۵۶	۲/۱۷۹۸	۰/۷۱۱۸	۲/۲۹۹۸	۰/۶۵۴۳	۲/۳۰۷۲	۳	۲

توجه اینکه با افزایش مقدار پارامتر وابستگی، دقت برآوردهای β کاهش پیدا می‌کند. به عبارت دیگر هنگامی که وابستگی فضایی مشاهدات افزایش پیدا می‌کند، دقت برآورد حاصل از مدل حاشیه‌ای برای پارامتر رگرسیونی کاهش پیدا می‌کند. نتایج شبیه‌سازی در حالت نامعلوم بودن پارامتر وابستگی فضایی، به دلیل وجود اشکال در برنامه محاسبه برآوردها ارائه نشده‌اند.

۵ بحث و نتیجه‌گیری

در مدل‌بندی داده‌های بقای وابسته فضایی، علاوه بر برآورد میزان تأثیر متغیرهای تبیینی مختلف بر تابع خطر، برآورد میزان وابستگی داده‌های بقا نیز مورد توجه است. در این مقاله، یک روش برآورد دو مرحله‌ای برای این منظور معرفی شد که بر اساس آن میزان تأثیر متغیرهای تبیینی بر تابع خطر، با استفاده از مدل حاشیه‌ای خطرهای متناسب برآورد می‌شود و از تابع مفصل زوجی کلاسیون برای لحاظ کردن وابستگی داده‌ها و سپس برآورد پارامتر وابستگی استفاده می‌شود. این روش محدودیتی در نوع تابع کوواریانس فضایی مشاهدات ندارد. کافی است تابع درست‌نمایی مرکب (۲.۳) نسبت به پارامترهای تابع کوواریانس، ماکسیمم شود. نتایج مطالعه شبیه‌سازی، بیانگر عملکرد مناسب این روش در برآورد ضرایب رگرسیونی و پارامتر وابستگی بود.

مراجع

- امیدی، م.، محمدزاده، م. (۱۳۹۰)، مدل‌بندی درست‌نمایی و بیز تجربی خشکسالی با استفاده از تابع مفصل، نشریه علوم دانشگاه خوارزمی، ۱۰، ۶۴۳-۶۵۴.
- مترجم، ک.، محمدزاده، م. و آبیاری، آ. (۱۳۹۴)، مدل‌بندی فضایی داده‌های بقای سانسور شده، پژوهش‌های ریاضی، ۱۰، ۶۱-۷۰.

Banerjee S. , Dey D. K. (2005), Semiparametric Proportional Odds Models for Spatially Correlated Survival Data, *Lifetime Data Analysis*, **11**, 175-191.

- Clayton D. (1978), A Model for Association in Bivariate Life Tables and Its Application in Epidemiological Studies of Familial Tendency in Chronic Disease Incidence, *Biometrika*, **65**, 141–151.
- Cox DR. (1975), Partial Likelihood, *Biometrika*, **62**, 269-276.
- Li Y. , Lin X. (2006), Semiparametric Normal Transformation Models for Spatially Correlated Survival Data, *Journal of the American Statistical Association*, **101**, 591–603.
- Li Y. , Ryan L. (2002), Modeling Spatial Survival Data Using Semiparametric Frailty Models, *Biometrics*, **58**, 287–297.
- Nelsen R.B. (2005), *An Introduction to Copulas*. Second Edition. Springer Series in Statistics .
- Pan C., Cai B., Wang L., Lin X. (2014), Bayesian SemiParametric Model for Spatial Interval-Censored Survival Data, *Computational Statistics and Data Analysis*, **74**, 198–209.
- Prenen L., Braekers R. (2017), Extending the Archimedean Copula Methodology to Model Multivariate Survival Data Grouped in Clusters of Variable Size, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, **79**, 483-505.
- Shiau J. T. (2006), Fitting Drought Duration and Severity with Two-Dimensional Copulas, *Water Resources Research*, **20**, 795-815.
- Zhou H., Hanson T., Knapp R. (2015), Marginal Bayesian Nonparametric Model for Time to Disease Arrival of Threatened Amphibian Populations, *Biometrics*, **71**, 1101-1110.



پیشگویی فضایی میدان تصادفی ناگوسی با استفاده از قضیه تصویر

اسحاق الماسی^۱، مهدی امیدی

گروه ریاضی، دانشگاه ایلام

چکیده: یکی از مهمترین مسائل آمار فضایی تعیین توزیع تحقق‌های میدان تصادفی به منظور پیشگویی در نقاط نامعلوم از میدان تصادفی است. فرض گاوسی بودن میدان باعث می‌شود استفاده از پیشگوهای خطی توجیه‌پذیر باشد. نامشخص بودن توزیع میدان تصادفی محققین را ملزم به استفاده از روش‌هایی می‌کند که بر اساس آن‌ها امکان پیشگویی ناگوسی میسر شود. در این مقاله یک روش ناپارامتری را بر اساس قضیه تصویر برای پیشگویی ارایه می‌دهیم و بر اساس آن مدلی برای پیشگویی در میدان‌های ناگوسی مطرح می‌شود. در ادامه با استفاده از مطالعه شبیه‌سازی میزان دقت و صحت این روش مورد بررسی قرار می‌گیرد.

واژه‌های کلیدی: آمار فضایی، پیشگویی ناگوسی، قضیه تصویر، آماره‌های مرتب.

کد موضوع بندی ریاضی (۲۰۱۰): 91B72، 62H11

۱ مقدمه

در آمار فضایی پیشگویی مقدار میدان تصادفی در مکان‌های فاقد مشاهده معمولاً با فرض گاوسی بودن میدان تصادفی انجام می‌شود. وجود این فرض باعث می‌شود که استفاده از پیشگوهای خطی مانند کریکینگ (کرسی، ۱۹۹۳) امکان‌پذیر باشد. در برخی شرایط ممکن است واقع‌گرایانه نباشد از این روش باید از روش‌هایی استفاده نمود که بتوان برای میدان‌های ناگوسی پیشگویی انجام داد. یکی از روش‌هایی که تاکنون در پیشگویی میدان‌های ناگوسی فضایی مطرح شده استفاده از تابع مفصل (امیدی و محمدزاده، ۲۰۱۸) است، بر اساس این روش با استفاده از محاسبات پیچیده توزیع میدان تصادفی به دست می‌آید اما محدودیت‌هایی از قبیل تعیین توزیع حاشیه‌ای و نسبی بودن تابع توزیع به دست آمده در آن مطرح است. براکول و داویس (۱۹۹۱) پیش بینی در سری‌های زمانی را با استفاده از قضیه تصویر مطرح کردند، الماسی و همکاران (۲۰۱۷) با استفاده از قضیه تصویر روشی را برای پیشگویی آماره‌های ترتیبی سانسور شده معرفی کردند. در این مقاله با استفاده از روش مطرح شده توسط الماسی و همکاران (۲۰۱۷) بر اساس نزدیکترین همسایه به نقاط گمشده فضایی ترکیب خطی

^۱ نام و ایمیل ارائه دهنده مقاله: اسحاق الماسی، i.almasi@ilam.ac.ir

ناپارامتری برای پیشگویی فضایی نقاط گمشده ارائه می‌شود. در ادامه با استفاده از شبیه‌سازی دقت و صحت این روش با استفاده از معیار RMSE مورد بررسی قرار می‌گیرد.

۲ ساختار همبستگی میدان تصادفی

در آمار فضایی مقدار میدان تصادفی $Z = \{Z(s); s \in D \subset R^d\}$ در موقعیت مشخص s بر اساس تحقیقات میدان تصادفی $s_1, \dots, s_n$ ، یعنی $z = (z(s_1), \dots, z(s_n))$ پیشگویی می‌شود. توابع میانگین و کواریانس فضایی میدان تصادفی به ترتیب به صورت

$$\begin{aligned}\mu(s) &= E(Z(s)); \quad (s) \in D \\ C(s, s') &= Cov(Z(s), Z(s')); \quad (s, s') \in D\end{aligned}$$

تعریف می‌شوند. میدان تصادفی فضایی-زمانی $Z(\cdot)$ مانای مرتبه دوم است هرگاه میانگین آن نسبت به (s) ثابت باشد، یعنی $\mu(s) = \mu$ و کواریانس فضایی-زمانی تابعی از فاصله موقعیت‌های فضایی

$$Cov(Z(s, s')) = C(\mathbf{h}_s); \quad (s, s') \in D$$

باشد، که در آن $\mathbf{h}_s = s - s'$ تأخیر فضایی نامیده می‌شود. اگر تابع کواریانس فقط تابعی از اندازه تأخیرها بوده و به جهت آنها بستگی نداشته باشد، همسانگرد نیز نامیده می‌شود.

تغییرنگار فضایی-زمانی به صورت $\gamma(s, s') = Var[Z(s) - Z(s')]$ تعریف می‌شود. چنانچه میانگین ثابت و تغییرنگار تابعی از تأخیرهای فضایی و زمانی باشد، یعنی $\gamma(\mathbf{h}_s) = Var[Z(s) - Z(s')]$ ، میدان مانای ذاتی نامیده می‌شود. تحت مانایی ضعیف تابع تغییرنگار به صورت

$$\gamma(\mathbf{h}_s) = \gamma[C(\cdot) - C(\mathbf{h}_s)]$$

با تابع کواریانس فضایی-زمانی در ارتباط است. توابع کواریانس فضایی ساختار همبستگی داده‌ها فضایی را تفسیر نموده و خواص آن‌ها به گونه‌ای است که همواره اندازه همبستگی مثبت را به دست می‌دهند و مقدار آن مشاهداتی با فاصله نزدیک حداکثر و برای فواصل دورتر مقدار آن کم می‌شود. **امیدی و محمدزاده (۲۰۱۶)** روش‌هایی برای ساخت توابع کواریانس همیشه مثبت ارائه دادند که به عنوان مثال تابع

$$C(\mathbf{h}_s) = \sigma^2 \exp(-\alpha \mathbf{h}_s^\beta) \quad (۱.۲)$$

برای σ و α مثبت و $\beta \in [0, 2]$ یک تابع کواریانس فضایی معتبر است.

۳ پیشگویی با قضیه تصویر

بر اساس قضیه تصویر بهترین پیش بینی خطی $Z = \{Z(s); s \in D \subset R^d\}$ بر اساس مشاهدات $s_1, \dots, s_n$ ، یعنی $z = (z(s_1), \dots, z(s_n))$ به صورت زیر

$$\hat{Z}(s_l) = \sum_{i=1}^n a_i z(s_i)$$

است، که در آن ضرایب ترکیب خطی از حل همزمان معادلات

$$\sum_{i=0}^n a_i E(Z(s_i)Z(s_j)) = E(Z(s_l)Z(s_j)), \quad j = 0, \dots, n$$

به دست می آید. اگر $i, k = 1, \dots, n$ فاصله اقلیدسی مشاهدات i -ام و k -ام و بردار $d_{(1)}, \dots, d_{(n)}$ مقادیر مرتب شده فواصل اقلیدسی $d_{il}, i = 1, \dots, n$ باشند، آنگاه بر اساس قضیه تصویر، بهترین پیشگوی خطی در مکان $z(s_l)$ به صورت

$$\hat{Z}(s_l) = a_0 + a_r z(s_r) + a_k z(s_k) \quad (1.3)$$

به دست می آید، که در آن a_0, a_r و a_k از طریق حل همزمان معادلات

$$\begin{cases} a_0 + a_r E(Z(s_r)) + a_k E(Z(s_k)) = E(Z(s_l)), \\ a_0 E(Z(s_r)) + a_r E(Z(s_r)^2) + a_k E(Z(s_r)Z(s_k)) = E(Z(s_r)Z(s_l)), \\ a_0 E(Z(s_k)) + a_r E(Z(s_r)Z(s_k)) + a_k E(Z(s_k)^2) = E(Z(s_l)Z(s_k)). \end{cases} \quad (2.3)$$

به دست می آید. با در نظر گرفتن فاصله از دو مشاهده از نزدیکترین مکان‌ها به نقطه فاقد مشاهده، ضرایب به صورت $a_0 = E(Z(s_l)) - a_r E(Z(s_r)) - a_k E(Z(s_k))$

$$a_r = \frac{\rho(d_{(1)}) - \rho(d_{kr})\rho(d_{(r)})}{1 - \rho^2(d_{kr})}, \quad a_k = \frac{\rho(d_{(r)}) - \rho(d_{kr})\rho(d_{(1)})}{1 - \rho^2(d_{kr})}$$

برآورد می شوند، که در آن $\rho(d_{kr}) = \text{Corr}(z(s_r), z(s_k))$ مقدار همبستگی فضایی میان $z(s_r)$ و $z(s_k)$ است.

۴ مطالعه شبیه سازی

در این بخش با استفاده از تابع کواریانس (۱.۲) یک میدان تصادفی فضایی به ازای تعداد مکان‌های مختلف تولید می شود. در ادامه برای مقادیر مختلف α و β صحت و دقت برآورد رابطه (۱.۳) در ۵۰، ۱۰۰، ۵۰۰ و ۱۰۰۰ مکان با استفاده از معیار کارایی نسبی RMSE مورد بررسی قرار می گیرد. به منظور کنترل اثر تصادفی بودن داده‌های تولید شده در دقت برآوردگر روند اجرای برنامه ده بار تکرار شده و میانگین RMSE به دست آمده در ۱۰ بار مورد بررسی قرار گرفت. در جدول ۱، میانگین RMSE را در ده بار تکرار مجزای پیشگویی فضایی نقاط تولید شده آمده است. هم‌طور که مشاهده می شود نمی توان اظهار داشت که افزایش تعداد نقاط و مقادیر α و β تاثیری در برآوردگر داشته است.

بحث و نتیجه گیری

در این مقاله، با استفاده از قضیه تصویر یک روش ناپارامتری برای پیش گویی در میدان‌های تصادفی ناگوسی مطرح شد. در ادامه با استفاده از مطالعه شبیه سازی میزان دقت و صحت این روش بر اساس معیار RMSE مورد بررسی قرار می گیرد. نتایج نشان داد تغییرات در مقادیر پارامترهای تابع کواریانس فضایی و همچنین افزایش تعداد مکان‌ها تاثیری بر دقت برآوردگر ندارد.

اگرچه در این مقاله برآوردگری برای پیشگویی ناگوسی مطرح شد اما این روش بر اساس نزدیکترین دو همسایه به دست آمد که خود می تواند از نقاط ضعف برآوردگر تلقی گردد. از این رو به دست آوردن برآوردگر برای همسایه‌های بیشتر و مقایسه این روش با پیشگوی کرکینگ و تابع مفصل فضایی از موضوعات مهم در مبحث پیشگویی فضایی است.

جدول ۱: میانگین RMSE پیشگویی‌ها

α					تعداد مکان‌ها
۲۰	۱۰	۵	۱	β	
۰/۹۳۲	۰/۹۴۱	۰/۹۴۵	۰/۸۸۸	۰/۵	۵۰
۰/۹۳۲	۰/۸۸۷	۰/۸۰۱	۰/۹۰۹	۱	
۰/۹۴۹	۰/۹۶۰	۰/۹۶۷	۱/۰۱۵	۱/۵	
۰/۹۴۹	۰/۹۹۲	۰/۸۹۵	۰/۸۶۸	۰/۵	۱۰۰
۰/۹۵۳	۰/۸۹۲	۰/۹۰۳	۱/۰۹۳	۱	
۰/۹۶۹	۰/۹۷۶	۱/۰۱۹	۰/۸۶۷	۱/۵	
۰/۹۹۱	۰/۹۸۴	۰/۹۱۹	۰/۹۸۷	۰/۵	۵۰۰
۰/۹۴۶	۱/۰۴۳	۱/۰۷۷	۰/۹۹۵	۱	
۱/۰۹۷	۱/۱۸۳	۱/۱۴۲	۱/۱۵۲	۱/۵	
۰/۹۹۷	۰/۹۷۹	۰/۹۳۵	۱/۱۴۲	۰/۵	۱۰۰۰
۰/۹۸۵	۱۱/۰۱	۱/۲۳۳	۱/۳۳۷	۱	
۱/۰۷۵	۱/۳۱۱	۱/۲۶۴	۰/۸۶۶	۱/۵	

مراجع

- Almasi, I., Mohammadpour, A., and Mohammadi, M. (2017), Best Linear Unbiased Interpolation of Order Statistics, *Communications in Statistics-Simulation and Computation*, **46**(5), 4161–4171.
- Arnold, B. C., Balakrishnan, N., and Nagaraja, H. N. (1992), *A First Course in Order Statistics* (Vol. 54), Siam.
- Brockwell, P. J., Davis, R. A. (1991). *Time Series: Theory and Methods*, Second ed., Springer Series in Statistics, Springer-Verlag, New York.
- Cressie, N. (1993), *Statistics for Spatial Data*, John Wiley, New York.
- Omidi, M. and Mohammadzadeh, M. (2016), A New Method to Build Spatio-Temporal Covariance Functions: Analysis of Ozone Data. *Statistical Papers*, **57**(3), 689-703.
- Omidi, M. and Mohammadzadeh, M. (2018), Spatial Interpolation Using Copula for non-Gaussian Modeling of Rainfall Data, *JIRSS*, **17**(2), 165-179.



رگرسیون توزیعی فضایی با دم‌های نیمه‌سنگین

جمیل اونق، حسین باغیشنی^۱

گروه آمار، دانشکده علوم ریاضی، دانشگاه صنعتی شاهرود

چکیده: در این مقاله، به مدل‌بندی داده‌های فضایی شبکه‌ای معطوف می‌شویم که در بین خود چند داده پرت دارند. برای این کار از رده جدید و بسیار منعطف مدل‌های رگرسیون توزیعی با در نظر گرفتن یک توزیع با دم نیمه‌سنگین برای پاسخ استفاده می‌کنیم. در این رده از مدل‌ها، همه پارامترهای خانواده توزیع پاسخ را می‌توان به متغیرهای تبیینی رگرسیونی و اثرات غیرخطی مثل اثر فضایی مرتبط کرد. برای مدل‌بندی اثر فضایی داده‌ها از یک میدان تصادفی مارکوفی گاوسی (GMRF) استفاده می‌کنیم و در یک مثال واقعی کارایی مدل پیشنهادی را نمایش می‌دهیم.

واژه‌های کلیدی: توزیع با دم نیمه‌سنگین، داده‌های شبکه‌ای، رگرسیون توزیعی، میدان تصادفی مارکوفی گاوسی.
کد موضوع‌بندی ریاضی (۲۰۱۰): 97K80، 62H11.

۱ مقدمه

بسیاری از داده‌ها به‌طور ذاتی دارای تغییرپذیری مکانی هستند که به آن‌ها داده‌های فضایی می‌گوییم. امروزه با پیشرفت‌های متعدد ثبت داده‌ها، اطلاعات فضایی داده‌ها به همراه کمیت‌های مختلف مد نظر ثبت می‌شوند که ذخیره‌ای با ارزش برای تحلیل و تصمیم‌سازی‌های کلان کشورها محسوب می‌شوند. مدل‌های رگرسیونی به‌عنوان پرکاربردترین مدل‌های آماری، وظیفه ایجاد ارتباط بین متغیر پاسخ هدف با متغیرهای تبیینی رگرسیونی را برعهده دارند. دو رده پرکاربرد رگرسیونی، مدل‌های خطی تعمیم‌یافته^۱ (GLMs) (نلدر و وودبرن، ۱۹۷۲) و مدل‌های جمعی تعمیم‌یافته^۲ (GAMs) (هیستی و تبشیرانی، ۱۹۹۰) هستند. در هر دوی این رده از مدل‌ها، توزیع متغیر پاسخ، Y ، عضوی از خانواده توزیع‌های نمایی در نظر گرفته می‌شود، که در آن پارامتر مکان توزیع متغیر پاسخ، μ ، به‌عنوان تابعی از متغیرهای تبیینی رگرسیونی مدل‌بندی می‌شود. یعنی در مدل‌های GLM و GAM سایر پارامترهای توزیع، مانند مقیاس و چولگی، به‌طور مستقیم با متغیرهای تبیینی تغییر نمی‌کنند.

^۱Generalized Linear Models

^۲Generalized Additive Models

^۱ نام و ایمیل ارائه دهنده مقاله: حسین باغیشنی، hbaghishani@shahroodut.ac.ir

مدل‌های GAM با در نظر گرفتن اثرات غیرخطی (ناپارامتری) متغیرهای تبیینی به صورت جمعی، نسبت به مدل‌های GLM، انعطاف بیشتری دارند، اما به غیر از پارامتر مکان اجازه مدل‌بندی سایر پارامترهای توزیع پاسخ را نمی‌دهند. برای رفع این محدودیت و مدل‌بندی توام پارامترهای مکان، مقیاس و شکل متغیر پاسخ، ریگی و استاسینوپولوس (۲۰۰۵) رده مدل‌های جمعی تعمیم‌یافته برای مکان، مقیاس و شکل^۳ (GAMLSS) را معرفی کردند. با توجه به آن‌که در این رده از مدل‌ها همه پارامترها را می‌توان به طور مستقیم به متغیرهای تبیینی و اثرات غیرخطی (ناپارامتری) مرتبط کرد، از عنوان رگرسیون توزیعی^۴ برای آن‌ها نیز استفاده کرده‌اند. به عنوان نمونه می‌توان **کلین و نیب** (۲۰۱۶) و **کلین و همکاران** (۲۰۱۵) را نام برد. یکی از قابلیت‌های مدل‌های رگرسیون توزیعی امکان افزودن اثرات تصادفی به هر کدام از پارامترهای توزیع شرطی پاسخ برای پاسخ‌های وابسته است. به عنوان مثال، می‌توان با استفاده از نظریه میدان‌های تصادفی مارکوفی^۵ (MRF)، یک اثر تصادفی فضایی را وارد مدل کرد (رو و هلد، ۲۰۰۵). بی‌سگ و کوپربرگ (۱۹۹۵) مدل اتورگرسیو شرطی ذاتی^۶ (ICAR) نرمال را که زیررده مهمی از MRF هستند معرفی و بررسی کردند. وود (۲۰۰۶) و رو و هلد (۲۰۰۵) مدل ICAR را در قالب رده GAM توسعه دادند. همچنین مدل ICAR توسط **دی‌بستانی و همکاران** (۲۰۱۸) برای مدل‌های رگرسیون توزیعی گسترش داده شد.

پیچیدگی‌های زندگی واقعی همیشه اثبات کرده‌اند که همه چیز به دلخواه و مطابق انتظار پیش نخواهند رفت. معمولاً در بین داده‌ها یک یا چند داده غیرمعمول نسبت به سایر داده‌ها حضور دارند. اگر چه حضور توزیع‌های مختلف چوله و دم سنگین در رده مدل‌های رگرسیون توزیعی، امکان مدل‌بندی داده‌های پرت را میسر می‌سازند، ولی در موقعیت‌هایی که تعداد آن‌ها نسبت به همه مشاهدات کوچک است، استفاده از توزیع‌های دم‌سنگین می‌تواند پیچیدگی بیش از اندازه وارد مساله کند و بر احتمال‌های مرکزی توزیع به طور نامطلوبی تاثیرگذار باشد. این واقعیت، ما را بر آن داشت تا از یک توزیع متقارن با دم نیمه‌سنگین برای مدل‌بندی داده‌های پرت در چارچوب رگرسیون توزیعی استفاده کنیم. برای این منظور، از توزیع هایپربولیک سکانت^۷ (HS) استفاده کردیم (فیشچر، ۲۰۱۳). بنابراین، هدف اصلی ما در این مقاله، معرفی یک مدل رگرسیون توزیعی (GAMLSS) فضایی جدید برای مدل‌بندی داده‌های فضایی با دماهای نیمه‌سنگین است.

در بخش دوم توزیع HS و چارچوب مدل‌های رگرسیون توزیعی را تشریح می‌کنیم. در بخش سوم ساخت اثر فضایی ICAR به عنوان یکی از اثرات تصادفی ممکن در این مدل‌ها را معرفی می‌کنیم. بخش چهارم شامل تحلیل مجموعه داده سرقت شهر کلمبوس در ایالت اوهایو آمریکا و مقایسه عملکرد سه مدل رقیب است. در پایان نیز نتیجه‌گیری خواهیم کرد.

۲ چارچوب مدل‌بندی رگرسیون توزیعی

اصالت توزیع HS به **فیشچر** (۱۹۲۱) برمی‌گردد. ویژگی‌های این خانواده از توزیع‌ها توسط **دود** (۱۹۲۵) و **پرک** (۱۹۳۲) مورد بررسی قرار گرفتند. توزیع HS حول پارامتر مکان، که میانگین توزیع است، متقارن است. **وگان** (۲۰۰۲) این توزیع را به حالت چوله تعمیم داد و **فیشچر و وگان** (۲۰۰۲) توزیع هایپربولیک سکانت تعمیم‌یافته چوله را معرفی کردند. **فیشچر** (۲۰۱۳) در کتاب خود تعمیم‌های مختلف این توزیع را گردآوری کرده است و با برآزش این خانواده به داده‌های متنوع، کاربردهای آن را نشان داده است.

³ Generalized Additive Models for Location, Scale, and shape

⁴ Distributional regression

⁵ Markov Random Fields

⁶ Intrinsic Conditional Autoregressive

⁷ Hyperbolic Secant

تابع چگالی توزیع HS به صورت زیر تعریف می‌شود:

$$f_Y(y; \mu, \sigma) = \frac{\gamma}{\pi\sigma} \frac{\exp\left(\frac{y-\mu}{\sigma}\right)}{\exp\left(\gamma \frac{y-\mu}{\sigma}\right) + 1}, \quad y \in \mathbb{R}$$

که در آن پارامتر μ میانگین توزیع و σ پارامتر مقیاس هستند. واریانس این توزیع $\frac{\pi^2}{6}\sigma^2$ است و برای نمایش توزیع از نماد $Y \sim HS(\mu, \sigma)$ استفاده می‌کنیم. این توزیع یک توزیع با دم نیمه‌سنگین محسوب می‌شود، به طوری که دم‌های آن بین دم‌های توزیع نرمال و توزیع کوشی (به عنوان یک توزیع با دم سنگین) است.

رگرسیون توزیعی یک چارچوب مدل‌بندی جامع و انعطاف‌پذیر برای متغیر پاسخ فراهم کرده است. توزیع‌های متقارن، چوله و دم‌سنگین (پیوسته و گسسته) متعددی برای توزیع متغیر پاسخ در بسته‌افزار gamlss در محیط نرم‌افزار R تعبیه شده‌اند. برای مشاهده فهرست این توزیع‌ها می‌توانید به [استاسینوپولوس و همکاران \(۲۰۱۷\)](#) مراجعه کنید. برخی از این توزیع‌ها تا چهار پارامتر را شامل می‌شوند که می‌توان آن‌ها را با μ, σ, ν و τ نشان داد، که به ترتیب بیانگر پارامترهای مکان (مثل میانگین)، مقیاس (مثل انحراف معیار) و شکل (مثل چولگی و کشیدگی) هستند. در چارچوب مدل‌های رگرسیون توزیعی می‌توان همه پارامترهای توزیع (شرطی) پاسخ را به صورت پارامتری یا ترکیبی جمعی از توابع ناپارامتری هموار از متغیرهای تبیینی و اثرات تصادفی پنهان مدل‌بندی کرد. همچنین می‌توان آمیخته‌ای از دو بخش پارامتری و ناپارامتری را به صورت مدلی نیمه‌پارامتری در نظر گرفت.

برای شروع فرض کنید مشاهدات y_i ، برای $i = 1, \dots, n$ ، به طور مستقل (مستقل شرطی) دارای تابع (چگالی) احتمال $f(y_i | \theta^i)$ با بردار پارامترهای $\theta^i = (\theta_{1i}, \theta_{2i}, \theta_{3i}, \theta_{4i})' = (\mu_i, \sigma_i, \nu_i, \tau_i)'$ باشند. بنا بر [ریگی و استاسینوپولوس \(۲۰۰۵\)](#) یک مدل رگرسیون توزیعی به صورت زیر تعریف می‌شود:

$$g_k(\theta_k) = X_k \beta_k + \sum_{j=1}^{J_k} h_{jk}(x_{jk}) \quad (۱.۲)$$

که در آن $g_k(\cdot)$ ، برای $k = 1, 2, 3, 4$ ، یک تابع پیوند معلوم بکنوا است که پارامتر θ_k را به متغیرهای تبیینی مرتبط می‌کند. همچنین، برای هر k و $j = 1, 2, \dots, J_k$ ، بردار پارامترهای J'_k بعدی، $\beta_k = (\beta_{1k}, \beta_{2k}, \dots, \beta_{J'_k})'$ ، ماتریس طرح معلوم $n \times J'_k$ بعدی، $h_{jk}(\cdot)$ یک تابع ناپارامتری هموار از متغیر تبیینی X_{jk} و x_{jk} ها بردارهایی با طول n از مشاهدات متغیرهای تبیینی با اثرات غیرخطی هستند. مدل (۱.۲) یک مدل نیمه‌پارامتری را در چارچوب مدل‌های GAM برای پارامتر θ_k تشکیل می‌دهد.

در مدل (۱.۲) اثرات هموار متغیرها را می‌توان به صورت یک اثر تصادفی، مشابه $h(x) = Z\gamma$ ، نمایش داد که در آن Z یک ماتریس پایه است که بر اساس مقادیر x ساخته می‌شود و γ برداری از ضرایب (اثرات) تصادفی است که فرض می‌شود دارای توزیع نرمال چندمتغیره است. با این نمایش، واضح است که می‌توان اثرات تصادفی را نیز به پیشگوی (۱.۲)، برای هر پارامتر، به صورت جمعی اضافه کرد. بنابراین مدل‌بندی داده‌های وابسته مثل داده‌های طولی، سری زمانی و فضایی در چارچوب مدل‌های رگرسیون توزیعی نیز ممکن می‌شود. با این مقدمه، مدل (۱.۲) را برای هر کدام از پارامترهای مکان، مقیاس و شکل می‌توان به صورت زیر نمایش داد:

$$g_1(\theta_k) = X_1 \beta_1 + \sum_{j=1}^{J_k} Z_{jk} \gamma_{jk}$$

که در آن Z_{jk} ماتریس پایه معلوم $n \times q_{jk}$ بعدی وابسته به متغیرهای تبیینی هستند و γ_{jk} بردار تصادفی q_{jk} بعدی با پذیره $\gamma_{jk} \sim N_{q_{jk}}(0, G_{jk}^{-1})$ است که در آن G_{jk}^{-1} معکوس (تعمیم‌یافته) $q_{jk} \times q_{jk}$ بعدی ماتریس متقارن G_{jk}

$G_{jk}(\lambda_{jk})$ است که ممکن است به ابرپارامتر λ_{jk} وابسته باشد. اگر ماتریس G_{jk} ناویژه باشد، بردار ضرایب γ_{jk} دارای یک توزیع ناسره است که چگالی آن متناسب است با $\exp\left(-\frac{1}{\gamma} \lambda_{jk} \gamma'_{jk} G_{jk} \gamma_{jk}\right)$. به‌عنوان مثال می‌توان به اثر تصادفی ICAR (رو و هلد، ۲۰۰۵) برای داده‌های فضایی شبکه‌ای اشاره کرد که یک توزیع ناسره دارد. تعیین ساختارهای مختلف برای ماتریس‌های Z و G انواع مختلفی از اثرات جمعی را نتیجه می‌دهد. این اثرات شامل مولفه‌های تصادفی، مولفه‌های هموار متغیرها، مولفه‌های سری زمانی و اثرات فضایی می‌شوند. برای مولفه‌های هموار انتخاب‌های مختلفی مانند اسپلاین‌های جریمه‌ای، اسپلاین‌های درجه سوم، و چندجمله‌ای‌های موضعی وجود دارند.

چارچوب مدل‌بندی رگرسیون توزیعی این امکان را فراهم می‌کند که توزیع شرطی متغیر پاسخ (به شرط اثرات تصادفی) هر توزیع دلخواهی چه از خانواده توزیع‌های نمایی چه خارج از آن باشد. این یک ویژگی خیلی منعطف و خوب محسوب می‌شود. البته به‌عنوان یک محدودیت توزیع اثرات تصادفی باید گاوسی باشد. در این مقاله، ما با تعبیه کردن توزیع HS به چارچوب مدل‌های رگرسیون توزیعی امکان مدل‌بندی منعطف داده‌های پیوسته‌ای را فراهم کرده‌ایم که یک یا چند داده پرت را شامل می‌شوند. مدل مورد نظر ما یک مدل نیمه‌پارامتری مکان-مقیاس است که در ساختار مدل (۱.۲) با دو پارامتر $\mu = \theta_1$ و $\sigma = \theta_2$ صدق می‌کند. همچنین برای در نظر گرفتن وابستگی فضایی داده‌های شبکه‌ای از یک اثر تصادفی ICAR استفاده می‌کنیم که در زیربخش بعدی آن را معرفی خواهیم کرد.

رهیافت استنباطی مد نظر برای مدل پیشنهادی، مبتنی بر درستنمایی است. با توجه به این‌که در این مدل از مولفه‌های ناپارامتری جمعی مختلف می‌توان استفاده کرد، برای جلوگیری از پیچیدگی بیش از حد، از یک چارچوب درستنمایی جریمه‌ای استفاده می‌شود. برای سادگی تشریح روش، فرض می‌کنیم داده‌ها مستقل باشند. بنابراین تابع لگاریتم درستنمایی مدل پیشنهادی به صورت

$$\ell(\mu, \sigma) = \sum_{i=1}^n \log(f(y_i; \mu_i, \sigma_i))$$

است. تابع درستنمایی جریمه‌ای برای این مدل به صورت زیر است:

$$\ell_p = \ell(\mu, \sigma) - \sum_{k=1}^2 \sum_{j=1}^{J_k} \gamma'_{kj} G_{kj}(\lambda_{kj}) \gamma_{kj}. \quad (2.2)$$

مجموعه پارامترهایی که باید برآورد شوند، عبارتند از پارامترهای بخش خطی مدل، $\beta = (\beta_1, \beta_2)'$ ، پارامترهای اثرات تصادفی، $\gamma = (\gamma_{11}, \dots, \gamma_{J_1 1}, \gamma_{12}, \dots, \gamma_{J_2 2})'$ ، و ابرپارامترهای مدل، $\lambda = (\lambda_{11}, \dots, \lambda_{J_1 1}, \lambda_{12}, \dots, \lambda_{J_2 2})'$. در این جا، پارامترهای خطی β و پارامترهای اثرات تصادفی γ ، به ازای مقادیر ثابتی برای ابرپارامترهای هموارساز λ ، با ماکسیم کردن تابع درستنمایی جریمه‌ای ℓ_p در (۲.۲) برآورد می‌شوند. برای این کار دو روش پایه‌ای با نام‌های الگوریتم‌های RS و CG (ریگبی و استاسینوپولوس، ۲۰۰۵) وجود دارند. هر دوی این روش‌ها از یک الگوریتم کمترین توان‌های دوم (جریمه‌ای) موزون تکراری برای محاسبه برآوردهای ماکسیم درستنمایی پارامترهای مذکور، با مفروض بودن ابرپارامترهای هموارساز، استفاده می‌کنند. روش RS تعمیم الگوریتم به کار گرفته شده توسط ریگبی و استاسینوپولوس (۱۹۹۶) است. روش CG نیز تعمیم الگوریتمی است که توسط کول و گرین (۱۹۹۲) معرفی شد. در الگوریتم CG مشتقات مرتبه اول تابع درستنمایی و مقادیر امید ریاضی مشتقات مرتبه دوم تابع درستنمایی نسبت به پارامترهای توزیع یعنی μ و σ مورد نیاز هستند. البته برای مشتقات مرتبه دوم، مشابه روش بهینه‌سازی BFGS (برویدن، ۱۹۷۶)، می‌توان از تقریب‌های عددی نیز

استفاده کرد. مشتقات اول توزیع HS نسبت به μ و σ عبارتند از

$$\frac{\partial \ell(\mu, \sigma)}{\partial \mu} = -\frac{1}{\sigma} + \frac{2}{\sigma} \frac{\exp\left(\frac{y - \mu}{\sigma}\right)}{\exp\left(\frac{y - \mu}{\sigma}\right) + 1}$$

$$\frac{\partial \ell(\mu, \sigma)}{\partial \sigma} = -\frac{1}{\sigma} - \frac{y - \mu}{\sigma^2} + \frac{2}{\sigma^2} \frac{y - \mu}{\exp\left(\frac{y - \mu}{\sigma}\right) + 1}.$$

برای مشتقات دوم نیز از تقریب‌های عددی استفاده کرده‌ایم. برای برآورد بردار ابرپارامترهای λ دو روش موضعی^۸ و فراموضعی^۹ (ریگبی و استاسینوپلوس، ۲۰۱۳) پیشنهاد شده‌اند. روش موضعی سریعتر و اجرای آن ساده‌تر است. در واقع در روش موضعی برآورد ابرپارامترها در هر مرحله از الگوریتم‌های RS یا CG انجام می‌شود، در حالی که در روش فراموضعی، از معیارهایی مثل اعتبارسنجی متقابل تعمیم‌یافته^{۱۰} (GCV) (وود، ۲۰۰۶)، معیار اطلاع آکاییک (AIC) (آکاییک، ۱۹۷۴)، معیار اطلاع بیزی شوارتز (BIC) (شوارتز، ۱۹۷۸)، و مقدار ماکسیمم درست‌نمایی برای برآورد ابرپارامترها استفاده می‌شود.

برای مقایسه و انتخاب مدل برتر در بین مدل‌های نامزد می‌توان از معیارهای AIC و BIC استفاده کرد. بررسی نیکویی برازش مدل‌ها نیز توسط تحلیل باقی‌مانده‌ها انجام می‌شود.

۳ اثر تصادفی GMRF

در تحلیل داده‌های ناحیه‌ای (مشبکه‌ای) می‌توان انتظار داشت که نواحی همسایه، در ناحیه تحت مطالعه، شباهت بیشتری نسبت به نواحی غیرهمسایه داشته باشند. در این حالت، معمولاً، ویژگی مارکوفی برای اثر تصادفی فضایی داده‌ها پذیرفته می‌شود؛ به این معنی که احتمال شرطی رخداد پیشامدی به شرط رخداد آن در نواحی همسایه از رخداد آن پیشامد در نواحی غیرهمسایه مستقل (شرطی) است. این یک ویژگی مارکوفی گرافی بدون جهت است که در آن هر رأس نشان‌دهنده ناحیه و وجود یال بین دو رأس نشان‌دهنده داشتن مرز مشترک (همسایگی) بین دو ناحیه است (رو و هلد، ۲۰۰۵). برای توضیح بیشتر فرض کنید $\mathcal{G} = (\nu, \epsilon)$ یک گراف بدون جهت باشد که شامل رأس‌های $\nu = \{1, 2, \dots, q\}$ و مجموعه یال‌های ϵ است که به صورت زوج‌های (m, t) هستند، به‌طوری که $m, t \in \nu$. با توجه به رو و هلد (۲۰۰۵) بردار تصادفی $\gamma = (\gamma_1, \dots, \gamma_q)^T$ به شرط گراف $\mathcal{G}$ یک GMRF با بردار میانگین μ و ماتریس دقت (معکوس ماتریس کوواریانس) متقارن λG است، اگر و فقط اگر تابع چگالی آن به صورت

$$\pi(\gamma) = \exp \left[-\frac{1}{2} \lambda (\gamma - \mu)' G (\gamma - \mu) \right] \quad (۱.۳)$$

باشد، که در آن $G_{mt} \neq 0$ اگر و فقط اگر، برای هر $(m, t) \in \epsilon$ ، $m \neq t$ ، در این جا، G_{mt} مولفه سطر m ام و ستون t ام ماتریس G است. مدل CAR که اولین بار توسط بی‌سگ (۱۹۷۴) معرفی شد و توسط بی‌سگ و کوپربرگ (۱۹۹۵)

^۸Locally

^۹Globally

^{۱۰}Generalized Cross Validation

توسعه داده شد، حالت خاصی از یک GMRF به شکل (۱.۳) است. مدل CAR را به صورت زیر نیز می توان تعریف کرد:

$$\gamma_i | \gamma_{-i} \sim N \left(\sum_j \alpha_{ij} \gamma_j, k_i \right)$$

که در آن $\gamma_{-i} = (\gamma_1, \gamma_1, \dots, \gamma_{i-1}, \gamma_{i+1}, \dots, \gamma_q)$ ، $k_i = \frac{1}{\lambda G_{ii}}$ ، $i = 1, 2, \dots, q$ و

$$\alpha_{ij} = \begin{cases} -\frac{G_{ij}}{G_{ii}} & i \neq j \\ . & i = j \end{cases}$$

با استفاده از لم بروک (۱۹۶۴)، می توان نشان داد که به ازای $\mu = 0$ برای تمام i و j ها، رابطه $\alpha_{ij} k_j = \alpha_{ji} k_i$ برقرار است که تضمین کننده متقارن بودن ماتریس G است (بی سگ و کوپریگ، ۱۹۹۵).

زیررده ای از GMRF نیز وجود دارد که در آن ماتریس G ناویژه است و به مدل ICAR معروف است. در واقع، این مدل یک نسخه حدی از مدل CAR است. مدل ICAR در مدل های خطی تعمیم یافته فضایی با پاسخ های شبکه ای، برای تبیین اثرات تصادفی که دارای ساختار فضایی هستند به کار برده می شوند (بنرجی و همکاران، ۲۰۱۴). دی بستانی و همکاران (۲۰۱۸) در رده مدل های رگرسیون توزیعی، ماتریس G را به شکل زیر گسترش دادند. فرض کنید W یک ماتریس همسایگی متقارن باشد که در آن مولفه های روی قطر اصلی صفر هستند و سایر مولفه های (i, j) در صورتی که دو ناحیه i و j ، با شرط $i \neq j$ ، مرز مشترک داشته باشند برابر ۱ و در غیر این صورت برابر صفر هستند. با این تعریف، ماتریس دقت G از رابطه $G = D_w - W$ ساخته می شود که در آن D_w ماتریس قطری است که اعضای قطر اصلی آن مجموع سطرهای ماتریس W است. با افزودن یک اثر ICAR به پیشگوی خطی مدل (۱.۲)، می توان وابستگی فضایی داده ها را وارد مدل پیشنهادی کرد. برازش مدل نیز با استفاده از تابع درستنمایی جریمه ای (۲.۲) به سادگی قابل انجام است که در آن جریمه بر حسب G ساخته شده توسط مدل ICAR تعریف می شود.

۴ داده های سرقت شهر کلمبوس

برای نمایش کاربرد مدل پیشنهادی ما، در این بخش از مجموعه داده حجم سرقت منزل در شهر کلمبوس ایالت اوهایو آمریکا استفاده کرده ایم. متغیر پاسخ حجم سرقت منزل و اتومبیل (به درصد در ۱۰۰۰ خانوار)، با نام crime، است. با توجه به متغیرهای تبیینی موجود در این مجموعه از داده ها، دو متغیر درآمد خانوار، با نام income، و قیمت منزل، با نام home.value، برای برازش مدل های رگرسیونی به کار گرفته شدند. با توجه به این که شهر کلمبوس دارای ۴۹ ناحیه تقسیم بندی شده است، ماتریس همسایگی مدل ICAR یک ماتریس متقارن با بعد ۴۹ است.

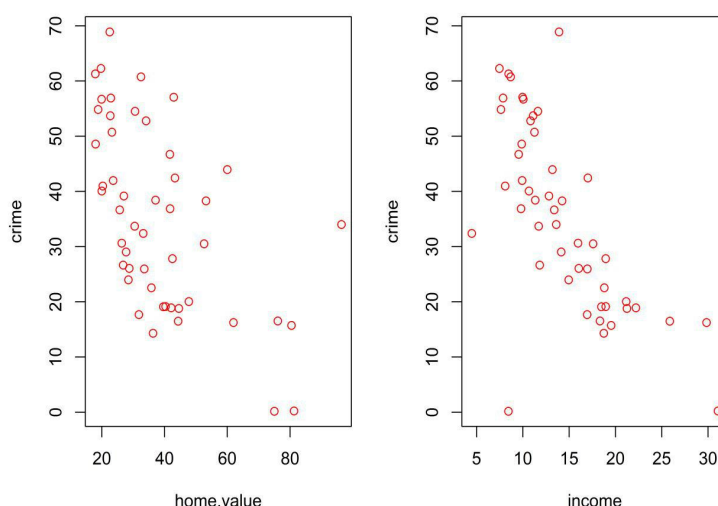
شکل ۱ نمودارهای پراکنش کناری متغیر پاسخ در مقابل متغیرهای تبیینی را نشان می دهد. همان طور که در دو نمودار پراکنش مشهود است، یک یا دو نقطه پرت حضور دارند که باید در مدل بندی لحاظ شوند. همچنین متغیر پاسخ به طور متوسط با افزایش سطوح متغیرهای تبیینی ارزش منزل مسکونی و درآمد روندی کاهشی دارد. در مقابل، با افزایش سطح درآمد تغییرات پاسخ نیز افزایش می یابد که می توان به مساله ناهم واریانس تا حدودی پی برد. با توجه به این تحلیل اکتشافی، از یک مدل مکان-مقیاس به صورت

$$\begin{cases} \mu = \beta_0 + \beta_1 income + h_1(home.value) + h_{spat}(district) \\ \sigma = \gamma_0 + \gamma_1 income \end{cases} \quad (1.4)$$

جدول ۱: نتایج برازش مدل‌ها برای مجموعه داده سرقت

مدل	$\hat{\beta}_0$	$\hat{\beta}_1$	$\hat{\gamma}_0$	$\hat{\gamma}_1$	$\hat{\nu}$	AIC	BIC
نرمال	۶۱/۱۶۹	-۱/۳۸۷	۳/۹۷۷	-۰/۱۲۹	-	۳۶۱/۶۳۹	۳۷۴/۸۸۱
HS	۵۴/۶۳۰	-۱/۳۳۳	۳/۷۲۳	-۰/۱۴۳	-	۳۵۷/۶۸۹	۳۷۰/۹۳۴
t-استیودنت	۶۰/۲۷۹	-۱/۳۳۲	۳/۹۵۰	-۰/۱۴۵	۱/۵۰۵	۳۶۰/۷۱۱	۳۷۵/۸۴۷

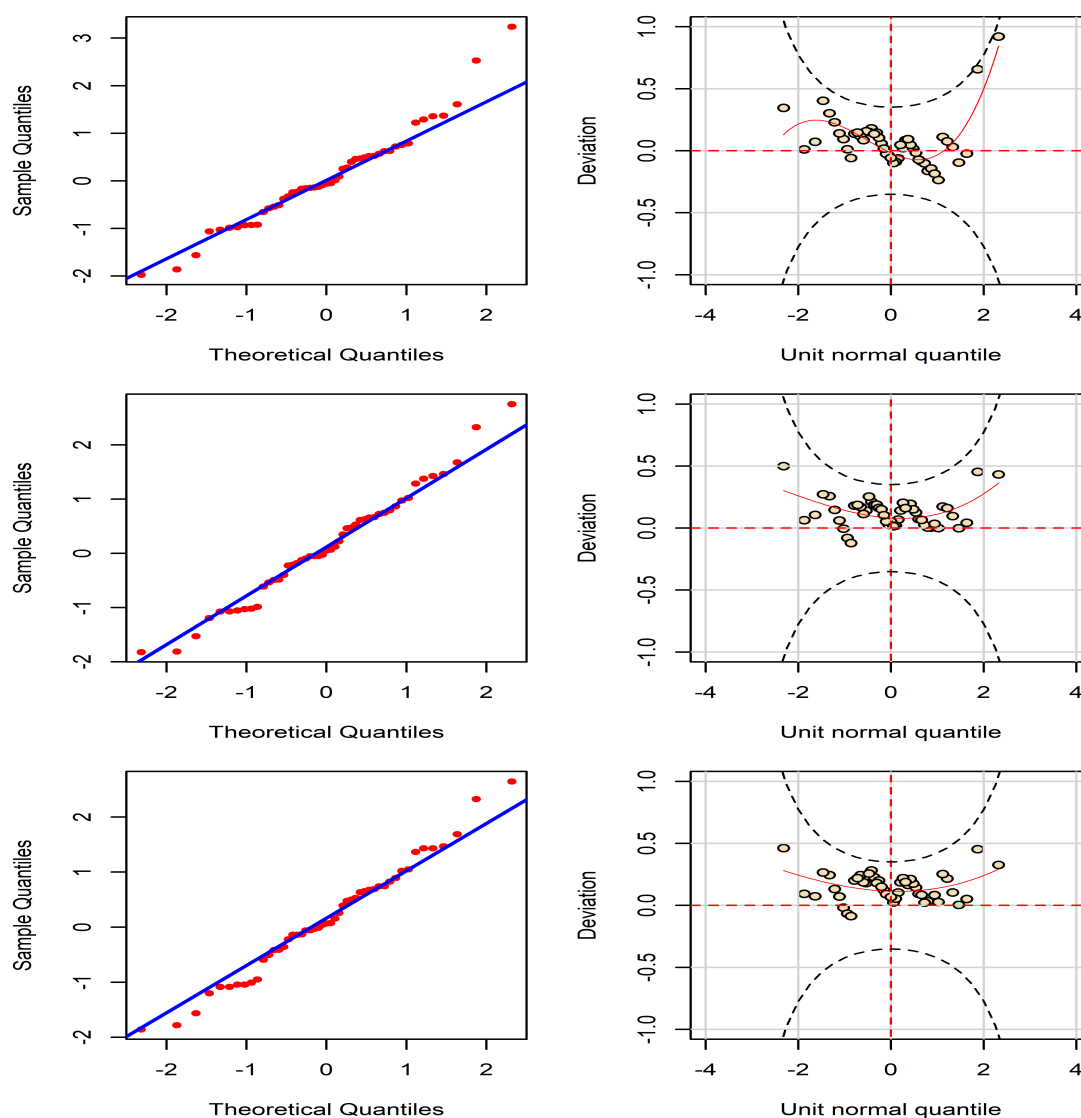
برای برازش داده‌ها استفاده کردیم. برای مقایسه عملکرد مدل HS، دو مدل رقیب مکان-مقیاس نرمال (به‌عنوان یک توزیع با دم سبک) و مدل مکان-مقیاس t-استیودنت (به‌عنوان یک توزیع دم سنگین) را نیز به داده‌ها برازش دادیم. نتایج برازش هر سه مدل در جدول ۱ گزارش شده‌اند. برآوردهای ضرایب رگرسیونی در هر دو پارامتر میانگین و مقیاس برای هر سه مدل تقریباً شبیه هم هستند. البته شباهت دو مدل HS و t-استیودنت بیشتر است. براساس هر دو معیار انتخاب مدل AIC و BIC، مدل برتر، مدل پیشنهادی ما است که مقادیر کوچکتری را نسبت به رقبای خود داراست. برای ارزیابی نیکویی برازش مدل‌ها، از نمودارهای چندک-چندک و کرمی^{۱۱} باقی‌مانده‌های روندزدایی شده استفاده کردیم. برای سه مدل برازش شده این نمودارها در شکل ۲ نمایش داده شده‌اند. با توجه به نمودارهای مدل نرمال می‌توان نقصان برازش آن را نتیجه گرفت. در مقابل برای هر دو مدل HS و t-استیودنت نیکویی برازش مورد تایید است.



شکل ۱: نمودارهای پراکنش کناری پاسخ و متغیرهای تبیینی برای داده‌های کلمپوس

با توجه به برتری مدل پیشنهادی ما، برآورد همه اثرات برای هر دو پارامتر میانگین و مقیاس را تنها برای مدل HS در شکل ۳ نشان داده‌ایم. جالب است که حجم سرقت در خانه‌های با ارزش کمتر یا با ارزش متوسط افزایش داشته است و برای خانه‌های گران‌قیمت به‌طور متوسط درصد سرقت کاهش یافته است. همچنین با افزایش درآمد خانوار، درصد سرقت کاهش می‌یابد. شکل ۴ نیز اثر فضایی برآورد شده را برای هر سه مدل نشان می‌دهد. اثر برآورد شده برای دو مدل HS و t-استیودنت شباهت زیادی به هم دارند و تفاوت‌هایی با مدل نرمال دارا می‌باشند. بنابراین با توجه به نقصان مدل نرمال، می‌توان از نتایج حاصل از آن صرف‌نظر کرد. با توجه به نتایج حاصل از مدل HS می‌توان نتیجه گرفت، درصد سرقت در نواحی مرکزی و شرقی، به‌ویژه جنوب شرقی شهر کلمپوس بیشتر از سایر نواحی است. همچنین کمترین درصد سرقت مربوط

¹¹Worm plot

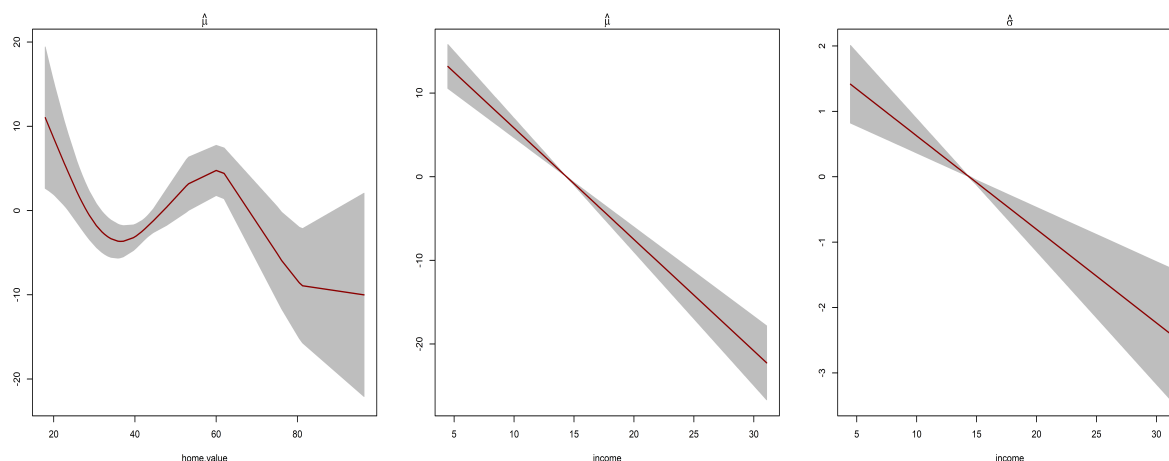


شکل ۲: نمودارهای چندک-چندک و کرمی مدل‌های نرمال (بالا)، HS (وسط) و t -استودنت (پایین)

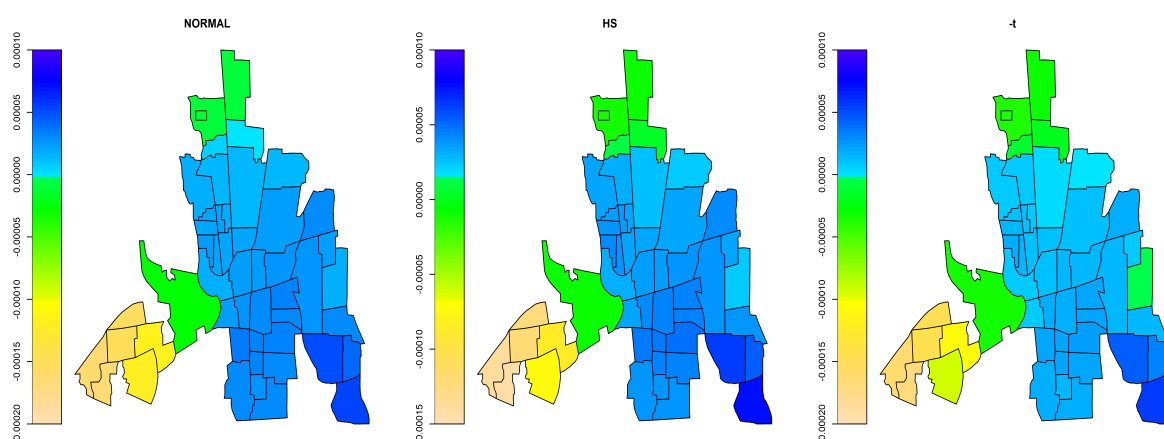
به نواحی جنوب غربی این شهر است.

بحث و نتیجه‌گیری

در این مقاله، یک مدل رگرسیون توزیعی فضایی جدید بر پایه توزیع با دم نیمه‌سنگین HS ارائه شد که می‌تواند جانشین مناسبی برای مدل نرمال در مواقعی باشد که بخش کوچکی از داده‌ها پرت محسوب می‌شوند. این مدل حتی در موقعیت‌هایی می‌تواند نسبت به مدل‌های تنومندی مثل مدل رگرسیون توزیعی t -استیودنت برتر باشد. با این کار، شاید حضور توزیع‌های با دم نیمه‌سنگین به‌عنوان جانشین‌هایی برای مدل‌های دم‌سنگین توجیه‌پذیر باشد. به‌ویژه مدل‌بندی داده‌های فضایی که به دلیل وابستگی فضایی داده‌ها تعداد داده‌های پرت چندان زیاد نیستند، می‌تواند نشانی از اهمیت ورود این دسته از مدل‌ها باشد.



شکل ۳: منحنی‌های اثرات برآوردشده پارامترهای خطی و غیرخطی مکان (وسط و سمت چپ) و مقیاس (سمت راست) برای مدل HS



شکل ۴: اثر فضایی برآوردشده پارامتر مکان در مدل‌های نرمال (سمت چپ)، HS (وسط) و t -استیودنت (سمت راست)

مراجع

Akaike, H. (1974), A New Look at the Statistical Model Identification, Selected Papers of Hirotugu Akaike, Springer, New York, 215-222.

Banerjee, S., Carlin, B. P., and Gelfand, A. E. (2014), *Hierarchical Modeling and Analysis for Spatial Data*, 2nd ed., Boca Raton: Chapman & Hall/CRC.

Besag, J., and Kooperberg, C. (1995), On Conditional and Intrinsic Autoregressions, *Biometrika*, **82**, 733–746.

- Besag, T. (1974), Spatial Interaction and the Statistical Analysis of Lattice Systems (With Discussion), *Journal of the Royal Statistical Society, Series B*, **36**, 192–236.
- Brook, D. (1964), On the Distinction Between the Conditional Probability and the Joint Probability Approaches in the Specification of Nearest-Neighbour Systems, *Biometrika*, **51**, 481–483.
- Broyden, C. G. (1976), Quasi-Newton Methods and Their Application to Function Minimization, *Mathematics of Computation*, **21**, 368–381.
- Cole, T. J., and Green, P. J. (1992), Smoothing Reference Centile Curves: The LMS Method and Penalized Likelihood, *Statistics in Medicine*, **11**, 1305–1319.
- De Bastiani, F., Rigby, R. A., Stasinopoulous, D. M., Cysneiros, A. H., and Uribe-Opazo, M. A. (2018), Gaussian Markov Random Field Spatial Models in GAMLSS, *Journal of Applied Statistics*, **45**(1), 168–186.
- Dodd, E. L. (1925), The Frequency Law of a Function of Variables with Given Frequency Laws, *Annals of Mathematics*, **27**, 12–20.
- Fischer, M. J. (2013), *Generalized Hyperbolic Secant Distributions: With Applications to Finance*, Springer Science & Business Media.
- Fischer, M., and Vaughan, D. C. (2002), Classes of Skewed Generalized Hyperbolic Secant Distributions, Working Paper No. 45 (Unpublished), Department of Statistics and Econometrics, FAU ErlangenNuremberg, Nuremberg.
- Fisher, R. A. (1921), 014: On the "Probable Error" of a Coefficient of Correlation Deduced From a Small Sample, *Metron*, **1**, 3–32.
- Hastie, T., and Tibshirani, R. (1990), Generalized Additive Models, *Statistical Science*, **1**, 297–310.
- Klein, N., and Kneib, T. (2016), Scale-Dependent Priors for Variance Parameters in Structured Additive Distributional Regression, *Bayesian Analysis*, **11**, 1071–1106.
- Klein, N., Kneib, T., Klasen, S., and Lang, S. (2015), Bayesian Structured Additive Distributional Regression for Multivariate Responses, *Journal of the Royal Statistical Society, Series C*, **64**, 569–591.
- Nelder, J. and Wedderburn, R. W. M. (1972), Generalized Linear Models, *Journal of the Royal Statistical Society: Series A (General)*, **135**, 370–384.
- Perks, W. (1932), On Some Experiments in the Graduation of Mortality Statistics, *Journal of the Institute of Actuaries*, **63**, 12–57.

- Rigby, R. A., and Stasinopoulos, D. M. (1996), A Semi-Parametric Additive Model for Variance Heterogeneity, *Statistics and Computing*, **6**, 57-65.
- Rigby, R. A., and Stasinopoulos, D. M. (2005), Generalized Additive Models for Location, Scale and Shape, *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, **54**, 507-554.
- Rigby, R. A., Stasinopoulos, M. D., and Voudouris, V. (2013), Discussion: A Comparison of GAMLSS with Quantile Regression, *Statistical Modelling*, **13**, 335-348.
- Rue, H., and Held. L. (2005), *Gaussian Markov Random Fields: Theory and Applications*, London: Chapman and Hall/CRC.
- Schwarz, G. (1978), Estimating the Dimension of a Model, *The Annals of Statistics*, **6**, 461-464.
- Stasinopoulos, M. D., Rigby, R. A., Heller, G. Z., Voudouris, V., and De Bastiani, F. (2017), *Flexible Regression and Smoothing Using GAMLSS in R*, London: Chapman and Hall/CRC.
- Vaughan, D. C. (2002), The Generalized Secant Hyperbolic Distribution and Its Properties, *Communications in Statistics-Theory and Methods*, **31**, 219-238.
- Wood, S. N. (2006), *Generalized Additive Models: An Introduction with R*, New York: Chapman & Hall.



تحلیل مدل‌های سری زمانی شمارشی فضایی

سپیده تشریفی^۱، امید کریمی، فاطمه حسینی
گروه آمار، دانشگاه سمنان

چکیده: داده‌های شمارشی چند متغیره برحسب زمان و مکان در خیلی از زمینه‌های کاربردی مشاهده می‌شوند. در این مقاله مدل‌های خطی سری زمانی شمارشی پواسنی روی این نوع از داده‌ها مورد مطالعه قرار می‌گیرند و همچنین همبستگی مکانی داده‌ها به وسیله یک مدل سری زمانی شمارشی فضایی با رهیافت بیزی را مورد تحلیل قرار می‌دهیم. در انتها یک مطالعه شبیه‌سازی برای بررسی مدل‌های معرفی شده، ارائه می‌شود.

واژه‌های کلیدی: سری‌های زمانی شمارشی، فرایند پواسن فضایی، داده‌های فضایی.
کد موضوع بندی ریاضی (۲۰۱۰): 62H11, 62M10, 62M30

۱ مقدمه

سری‌های زمانی مجموعه‌ای از مشاهدات هستند که بر حسب زمان مرتب شده است. ایجاد و به کارگیری مدل‌های آماری و تصادفی در قالب تحلیل سری زمانی، امروزه به کمک رایانه‌های پرسرعت، بسیار فراگیر شده و حاصل آن مدل‌هایی است که با داشتن پارامترهای بسیار انعطاف‌پذیر، می‌توانند آینده را برای هر پدیده‌ای پیش‌بینی کنند. کاربردهای تحلیل سری‌های زمانی، در زمینه‌های مختلف نظیر، کسب و کار، امور مالی، بورس، مهندسی و غیره دیده می‌شود. به بیان دیگر تحلیل سری زمانی، ایجاد مدلی گذشته‌نگر است تا امکان تصمیمات آینده‌نگر را فراهم سازد. معمولاً در سری زمان-گسسته، داده‌ها در مقاطع مشخصی از زمان مثل ساعت، روز یا هفته و حتی سال جمع‌آوری می‌شوند. در حالتی که داده‌ها از یک فرایند شمارشی به دست آمده باشد؛ سری زمانی حاصل را سری زمانی شمارشی می‌گویند (زهو و جو، ۲۰۰۶؛ بروکول و داویس، ۱۹۹۱). مدل‌های گوناگونی برای سری زمانی شمارشی وجود دارد. موردی که در مدل‌های سری زمانی شمارشی باید در نظر گرفته شود این است که وابستگی میان مشاهدات به‌طور مناسب مدل شود. یکی از این شیوه‌ها به کارگیری مدل‌های خطی تعمیم یافته است (مک‌کلاخ و نلدر، ۱۹۸۱).

^۱ نام و ایمیل ارائه دهنده مقاله: سپیده تشریفی، s.tasharofy@gmail.com

در این مقاله ابتدا مدل‌های سری زمانی برای تحلیل داده‌های شمارشی را مورد مطالعه قرار می‌دهیم، سپس همبستگی مکانی به وسیله تلفیق یک مدل پواسن فضایی با مدل‌های سری زمانی شمارشی مورد تحلیل قرار می‌دهیم و برخی استنباط‌های آماری مرتبط را بررسی می‌کنیم. تعدادی از این مدل‌ها را می‌توان در زمینه‌های کاربردی اقتصاد سنجی، علوم اجتماعی و فیزیک استفاده کرد. برای تحلیل رگرسیونی داده‌های شمارشی، مدل خطی معمولی مناسب نیست. زیرا که متغیر پاسخ مقادیر گسسته می‌گیرد. یک مدل مناسب، مدل رگرسیون پواسن است و یک نقطه شروع طبیعی برای گسترش آن به داده‌های شمارشی وابسته است. مدل رگرسیون پواسن در بسیاری از زمینه‌های کاربردی مورد استفاده قرار گرفته است. در واقع مدل پواسن ابزار اصلی برای مدل‌سازی داده‌های سری زمانی شمارشی است. با این حال، دیگر توزیع‌ها ممکن است به جای آن استفاده شود. بهترین آن در میان توزیع‌ها، توزیع دوجمله‌ای منفی است. صرف‌نظر از توزیع انتخاب شده عمدتاً مدل‌هایی را مورد بررسی قرار خواهیم داد که در چارچوب مدل‌های خطی تعمیم یافته سری زمانی است. این کلاس مدل‌ها و قضیه‌ی ماکسیمم درست‌نمایی یک چارچوب سیستماتیک برای تجزیه و تحلیل داده‌های سری زمانی به صورت کمی و کیفی ارائه می‌دهند. در واقع برآورد، تشخیص، ارزیابی مدل و پیش‌بینی به روش ساده انجام می‌شود. کلاس‌های جایگزین دیگری از مدل‌های رگرسیون وجود دارد. یکی از برجسته‌ترین آن‌ها مدل اتورگرسیون صحیح مقدار است. این مدل‌ها مشابه ساختار فرآیند اتورگرسیون معمول است.

ساختار مقاله به این صورت است که در بخش ۲ مدل‌های سری زمانی شمارشی معرفی و مورد مطالعه قرار می‌گیرند. همچنین جزئیات مدل‌های خطی و لگاریتم خطی سری زمانی شمارشی بیان می‌شود. در بخش ۳ یک مدل اتورگرسیون همبسته فضایی برای داده‌های شمارشی بیان و در بخش ۴ یک مطالعه شبیه‌سازی با مدل‌های معرفی شده، مورد مقایسه قرار می‌گیرد. در انتها بحث و نتیجه‌گیری ارائه می‌شود.

۲ مدل‌های سری زمانی شمارشی

۱.۲ مدل‌سازی رگرسیون پواسن

توزیع پواسن معمولاً برای مدل نرخ رویدادهای تصادفی استفاده می‌شود که در برخی از فاصله‌های زمانی ثابت رخ می‌دهد. اگر فرض کنیم λ نشان دهنده نرخ ورودی است؛ توزیع متغیر تصادفی Y نشان دهنده نرخ ورودی‌ها در یک فاصله زمانی ثابت است و دارای توزیع پواسن با تابع چگالی احتمال زیر است:

$$P[Y = y] = \frac{\exp(-\lambda)\lambda^y}{y!}, \quad y = 0, 1, 2, \dots \quad (1.2)$$

میانگین و واریانس Y هر دو برابر λ است. در اکثر مسائل کاربردی، داده‌های شمارشی معمولاً با برخی اطلاعات متغیر کمی مشاهده می‌شود (مک‌کلاخ و نلدر، ۱۹۸۱). به طور کلی فرض کنید که $X_1, \dots, X_p$ ، متغیر رگرسیونی مشاهده شده همراه با متغیر پاسخ شمارشی Y با توزیع پواسن است. در آن صورت مدل رگرسیونی به صورت $\lambda = \beta_0 + \sum_{i=1}^p \beta_i X_i$ است. این مدل چند مشکل برای برآورد دارد، زیرا که پارامتر λ باید مثبت باشد. یک انتخاب ساده برای مدل‌سازی رگرسیون داده‌های شمارشی، مدل لگاریتم خطی به صورت $\log \lambda = \beta_0 + \sum_{i=1}^p \beta_i X_i$ است. هر دو مدل متعلق به کلاس مدل‌های خطی تعمیم یافته‌اند که توسط نلدر و ویدربرن (۱۹۷۲) معرفی شده است. در تعریف مدل‌های سری زمانی شمارشی از مدل اتورگرسیون مرتبه یک ($AR(1)$) به صورت

$$Y_t = b_1 Y_{t-1} + \epsilon_t, \quad (2.2)$$

استفاده می‌شود. که در آن $|b_1| < 1$ و $\{\epsilon_t\}$ دنباله متغیرهای تصادفی نرمال با میانگین صفر و واریانس σ^2 است. این یک مدل استاندارد برای تجزیه و تحلیل سری‌های زمانی مقادیر واقعی است که مقدار فرآیند در زمان t بسته به مقدار فرآیند در زمان $t-1$ به همراه یک خطای تصادفی است. برای مثال [پریستی \(۱۹۸۱\)](#) را ببینید.

۱.۱.۲ مدل‌های خطی برای سری‌های زمانی شمارشی

فرض کنید که $\{Y_t\}$ سری‌های زمانی شمارشی را نشان می‌دهد. مدل اتورگرسیو پواسون را برای فرآیند پاسخ $\{Y_t\}$ به صورت

$$Y_t | \mathcal{F}_{t-1} \sim \text{Poisson}(\lambda_t), \quad \lambda_t = d + b_1 Y_{t-1}, \quad t \geq 1, \quad (3.2)$$

تعریف می‌کنیم که در آن d, b_1 و $\mathcal{F}_t = \sigma(Y_s, s \leq t)$ پارامترهای نامنفی و $\{\lambda_t\}$ به‌عنوان میانگین فرآیند Y_t با توجه به گذشته آن است. مقادیر مثبت d و b_1 ما را مطمئن می‌سازد که $\lambda_t > 0$ است زیرا که Y_t یک عدد صحیح نامنفی است. با توجه به این نکات، واضح است که مدل (۳.۲) با مولفه تصادفی توزیع پواسون و مؤلفه سیستماتیک $\eta_t = d + b_1 Y_{t-1}$ و تابع ربط همانی در چارچوب مدل‌های خطی تعمیم یافته قرار می‌گیرد. مدل (۳.۲) همان پویایی مدل (۲.۲) را نشان می‌دهد، بنابراین داریم:

$$Y_t = \lambda_t + (Y_t - \lambda_t) = d + b_1 Y_{t-1} + \epsilon_t \quad t \geq 1 \quad (4.2)$$

که در آن $\{\epsilon_t\}$ دنباله نوفه سفید است. این یک دنباله‌ای از متغیرهای تصادفی ناهمبسته با میانگین صفر و واریانس ثابت است. حالت کلی مدل (۳.۲) با مرتبه‌ی (p, q) به صورت

$$Y_t | \mathcal{F}_{t-1}^{Y, \lambda} \sim \text{Poisson}(\lambda_t), \quad \lambda_t = d + \sum_{i=1}^p a_i \lambda_{t-i} + \sum_{j=1}^q b_j Y_{t-j}, \quad t \geq \max(p, q) \quad (5.2)$$

تعریف می‌شود که در آن $\mathcal{F}_t^{Y, \lambda}$ یک σ میدان تشکیل شده از $\{Y_0, \dots, Y_t, \lambda_0\}$ است. یعنی $\mathcal{F}_t^{Y, \lambda} = \sigma(Y_s, \lambda_0, s \leq t)$ و $\{\lambda_t\}$ همانند قبل یک فرآیند شدت پواسون است. مانایی مرتبه دوم با رابطه $1 < \sum_{j=1}^p a_j + \sum_{j=1}^q b_j < \infty$ فراهم می‌شود. برای اطلاعات بیشتر به منابع [فرلند و همکاران \(۲۰۰۶\)](#) [ریدبرگ و شفر \(۲۰۰۰\)](#) [استریت \(۲۰۰۰\)](#) [هنین \(۲۰۰۳\)](#) مراجعه کنید. برای توزیع پواسون، $E[Y_t | \mathcal{F}_{t-1}^{Y, \lambda}] = \text{Var}[Y_t | \mathcal{F}_{t-1}^{Y, \lambda}] = \lambda_t$ را داریم. علاوه بر این، مدل (۳.۲) را می‌توان به عنوان واریانس ناهمسانی شرطی اتورگرسیو صحیح مقدار $INGARCH(1, 1)$ تعریف کرد. یعنی یک مدل واریانس ناهمسانی شرطی اتورگرسیو $GARCH$ صحیح مقدار. زیرا که ساختار آن مشابه مدل $GARCH$ معمول است که به موجب آن نوسانات بر روی مقادیر گذشته‌ی خود و مجذور پاسخ رگرسیون می‌شود. با در نظر گرفتن تابع پیوند کانونی $\nu_t \equiv \log \lambda_t$ خانواده مدل‌های لگاریتم خطی اتورگرسیو به صورت

$$Y_t | \mathcal{F}_{t-1}^{Y, \nu} \sim \text{Poisson}(\lambda_t), \quad \nu_t = d + a_1 \nu_{t-1} + b_1 \log(Y_{t-1} + 1), \quad t \geq 1. \quad (6.2)$$

حاصل می‌شود که در آن $\mathcal{F}_t^{Y, \nu}$ یک σ میدان تولید شده از $\{Y_0, \dots, Y_t, \nu_0\}$ است یعنی $\mathcal{F}_t^{Y, \nu} = \sigma(Y_s, \nu_0, s \leq t)$ به‌طور کلی، پارامترهای d, a_1, b_1 می‌توانند مثبت یا منفی باشند. اما باید شرایط خاصی برقرار باشد تا یک سری زمانی مانا به‌دست آوریم.

¹ Integer-Valued Autoregressive conditional heteroskedasticity

² Autoregressive conditional heteroskedasticity

۲.۱.۲ تحلیل درستنمایی

در این بخش تحلیل ماکسیمم درستنمایی شرطی را برای مدل خطی (۳.۲) به طور مختصر بیان می‌کنیم. این روش برای مدل‌های (۵.۲) و (۶.۲) کاملاً مشابه است. در مدل (۳.۲) پارامتر θ را به صورت یک بردار سه بعدی با پارامترهای مجهول $\theta = (d, a_1, b_1)'$ در نظر بگیرید. مقدار اولیه پارامتر به صورت $\theta_0 = (d_0, a_{10}, b_{10})'$ است. سپس تابع درستنمایی شرطی برای θ بر اساس مدل (۳.۲) با توجه به مقدار اولیه λ_0 توسط رابطه‌ی زیر به دست می‌آید:

$$L(\theta) = \prod_{t=1}^n \frac{\exp(-\lambda_t(\theta)) \lambda_t^{Y_t}(\theta)}{Y_t!}.$$

که در آن از فرض پواسون $\lambda_t(\theta) = d + a_1 \lambda_{t-1}(\theta) + b_1 Y_{t-1}$ توسط (۳.۲) و $\lambda_t = \lambda_t(\theta_0)$ استفاده می‌کنیم. از این رو، لگاریتم تابع درستنمایی را می‌توان به صورت

$$\ell(\theta) = \sum_{t=1}^n \ell_t(\theta) = \sum_{t=1}^n (Y_t \log(\lambda_t(\theta)) - \lambda_t(\theta)), \quad (۷.۲)$$

نوشت. همچنین تابع امتیاز به صورت زیر به دست می‌آید:

$$S_n(\theta) = \frac{\partial \ell(\theta)}{\partial \theta} = \sum_{t=1}^n \frac{\partial \ell_t(\theta)}{\partial \theta} = \sum_{t=1}^n \left(\frac{Y_t}{\lambda_t(\theta)} - 1 \right) \frac{\partial \lambda_t(\theta)}{\partial \theta}, \quad (۸.۲)$$

که در آن $\partial \lambda_t(\theta) / \partial \theta$ یک بردار سه بعدی با اجزای داده شده در رابطه‌ی زیر است:

$$\frac{\partial \lambda_t}{\partial d} = 1 + a_1 \frac{\partial \lambda_{t-1}}{\partial d}, \quad \frac{\partial \lambda_t}{\partial a_1} = \lambda_{t-1} + a_1 \frac{\partial \lambda_{t-1}}{\partial a_1}, \quad \frac{\partial \lambda_t}{\partial b_1} = Y_{t-1} + a_1 \frac{\partial \lambda_{t-1}}{\partial b_1}. \quad (۹.۲)$$

اگر جواب معادله $S_n(\theta) = 0$ ، وجود داشته باشد؛ برآوردگر ماکسیمم درستنمایی شرطی θ را به دست می‌دهد که توسط $\hat{\theta}$ مشخص می‌شود. علاوه بر این، همان‌طور که در رابطه زیر مشاهده می‌شود؛ ماتریس هسین در مدل (۳.۲) با تفاوت بیشتری از معادله‌ی (۸.۲) به صورت

$$\begin{aligned} H_n(\theta) &= - \sum_{t=1}^n \frac{\partial^2 \ell_t(\theta)}{\partial \theta \partial \theta'}, \\ &= \sum_{t=1}^n \frac{Y_t}{\lambda_t^2(\theta)} \left(\frac{\partial \lambda_t(\theta)}{\partial \theta} \right) \left(\frac{\partial \lambda_t(\theta)}{\partial \theta} \right)' - \sum_{t=1}^n \left(\frac{Y_t}{\lambda_t(\theta)} - 1 \right) \frac{\partial^2 \lambda_t(\theta)}{\partial \theta \partial \theta'}, \end{aligned} \quad (۱۰.۲)$$

به دست می‌آید. ماتریس اطلاع شرطی توسط رابطه زیر به دست می‌آید:

$$G_n(\theta) = \sum_{t=1}^n Var \left[\frac{\partial \ell_t(\theta)}{\partial \theta} \middle| \mathcal{F}_{t-1}^{Y, \lambda} \right] = \sum_{t=1}^n \frac{1}{\lambda_t(\theta)} \left(\frac{\partial \lambda_t(\theta)}{\partial \theta} \right) \left(\frac{\partial \lambda_t(\theta)}{\partial \theta} \right)', \quad (۱۱.۲)$$

این رابطه نقش مهمی در توزیع مجانبی برآورد ماکسیمم درستنمایی $\hat{\theta}^{MLE}$ دارد. به طور خاص تحت شرایط نظم، می‌توان ثابت کرد که $\hat{\theta}$ سازگار و مجانباً نرمال است. یعنی: $\sqrt{n}(\hat{\theta} - \theta_0) \rightarrow \mathcal{N}(0, G^{-1})$ ، که در آن G به صورت زیر تعریف می‌شود:

$$G(\theta) = E \left[\frac{1}{\lambda_t} \left(\frac{\partial \lambda_t}{\partial \theta} \right) \left(\frac{\partial \lambda_t}{\partial \theta} \right)' \right],$$

که در آن $E[\cdot]$ امید ریاضی است. تمام مقادیر فوق را می‌توان محاسبه کرد و برای پیش‌بینی‌ها، فواصل اطمینان و غیره مورد استفاده قرار می‌گیرد. اگرچه فرمول‌های بالا برای مدل خطی (۳.۲) ارائه شده‌اند. اما می‌توانند به‌طور مناسب برای مدل لگاریتم خطی (۶.۲) اصلاح شوند. در ادامه یک مدل سری زمانی شمارشی فضایی بیان و مورد مطالعه قرار می‌گیرد که تاثیر مکان را به این نوع مدل‌ها اضافه می‌کنیم.

۳ یک مدل اتورگرسیو همبسته فضایی برای داده‌های شمارشی

داده‌های شمارشی در آمار فضایی برخلاف سری زمانی عمدتاً از طریق یک فرآیند پنهان مورد بررسی قرار می‌گیرد. به عنوان مثال، مدل $INGARCH$ که توسط **فرلند و همکاران** (۲۰۰۶) و **هنین** (۲۰۰۳) بیان شده است و در آن مدل داده‌ها بواسون است و تابعی از شمارش‌های گذشته و امیدهای گذشته آن است.

ساختار مدل به این گونه است که، $Z(s_i, t)$ را مشاهدات در مکان‌های فضایی ثابت به صورت $s_i = \{s_1, \dots, s_n\}$ و نقاط زمانی $t = \{1, \dots, T\}$ در نظر می‌گیریم. سپس، $Y(s_i, t)$ را به عنوان یک متغیر تصادفی پنهان در موقعیت مکانی - زمانی (s_i, t) در یک ناحیه فضایی به صورت $s_j \in N_i$ یک همسایه فضایی s_i تعریف می‌کنیم. توجه داشته باشید که s_i به طور معمول یک بردار در $\mathbb{R}^2$ است. فرض کنید $\mathcal{H}_{Z(s_i, t)}$ گذشته فرایند مشاهدات در موقعیت مکانی s_i تا دوره‌ی زمانی t باشند. فرآیند واریانس ناهمسانی شرطی اتورگرسیو همبسته‌ی فضایی $SPINGARCH$ ، $[Z(s_i, t)|Y(s_i, t), \mathcal{H}_{Z(s_i, t)}]$ دارای توزیع بواسون شرطی با تابع احتمال

$$f(Z(s_i, t)|Y(s_i, t), \mathcal{H}_{Z(s_i, t)}) = \exp[A_i(Y(s_i, t), \mathcal{H}_{Z(s_i, t)})Z(s_i, t) - B_i(Y(s_i, t), \mathcal{H}_{Z(s_i, t)})C(Z(s_i, t))](1.3)$$

است، که در آن

$$\begin{aligned} \exp[A_i(Y(s_i, t), \mathcal{H}_{Z(s_i, t)})] &= \exp[Y(s_i, t)] + \eta Z(s_i, t-1) + \kappa E[Z(s_i, t-1)|Y(s_i, t-1), \mathcal{H}_{Z(s_i, t-1)}] \\ B_i(Y(s_i, t), \mathcal{H}_{Z(s_i, t)}) &= \exp[A_i(Y(s_i, t), \mathcal{H}_{Z(s_i, t)})] \\ C(Z(s_i, t)) &= -\log[Z(s_i, t)], \end{aligned}$$

به‌طوری‌که

$$\begin{aligned} Y(s_i, t)|y_t(N_i) &\sim N(\mu(s_i, t), \sigma^2) \\ \mu(s_i, t) &= \alpha(s_i) + \zeta \sum_{s_j \in N_i} \{y(s_j, t) - \alpha(s_j)\}. \end{aligned} \quad (2.3)$$

مدل $SPINGARCH$ از ترکیب (۱.۳)، که یک زنجیره مارکوف در زمان و (۲.۳) که یک زنجیره مارکوف در فضا است، تشکیل شده است. مشاهده $Z(s_i, t)$ روی فرایند گذشته کامل و فرآیند پنهان حال $Y(s_i, t)$ با استفاده از توزیع شرطی کامل در فضا در هر نقطه‌ی زمانی، مدل‌بندی شده است. فرضیات مارکوف، شرط را به یک مرحله زمانی و همسایگی‌های فضایی کاهش می‌دهد. قرار دهید $\mathbf{A}$ یک ماتریس همسایگی^۵ باشد به‌طوری‌که اگر s_i و s_j همسایه باشند؛ ورودی $(i, j) = 1$ را داریم و ζ را به $(\psi_{(1)}^{-1}, \psi_{(n)}^{-1}) \in \zeta$ محدود می‌کنیم که $\psi_{(k)}$ ، k امین بزرگترین مقدار ویژه‌ی $\mathbf{A}$ است. نتیجه‌ی مدل فرآیند پنهان در هر زمان t ، یک اتورگرسیو شرطی یا مدل CAR است که در **کریستی و ویکل** (۲۰۱۵) استفاده شده است. فرض کنید $\alpha = (\alpha(s_1), \dots, \alpha(s_n))$ باشد؛ آنگاه CAR دارای توزیع توام $\mathbf{Y}_t \sim N(\alpha, (I - C)^{-1}M)$ است که C یک ماتریس $n \times n$ با ورودی‌های ζ است و اگر $s_j, t \in N(s_i)$ باشد؛ ζ مقدار یک و در غیر این صورت برابر با ۰ خواهد بود. M یک ماتریس قطری با ورودی σ است. برای نمادگذاری راحت‌تر، $(s_i, t) \equiv \exp[A_i(\cdot)]$ را تعریف می‌کنیم. برای

⁴Spatially Correlated Integer Autoregressive Conditionally Heteroskedastic

⁵Neighborhood matrix

کرد. برای سادگی نمادها فرض کنید θ بردار پارامترهای σ^2 و ζ را نشان دهد و $\Sigma(\theta) \equiv (I - C)^{-1}M$ به عنوان ماتریس کوواریانس پنهان $n \times n$ باشد.

به طور خلاصه مدل *SPINGARCH* را می توان به صورت زیر نوشت:

$$\begin{aligned} Z(s_i, t) | \lambda(s_i, t) &\sim \text{pois}(\lambda(s_i, t)) \\ E[Z(s_i, t) | \lambda(s_i, t)] &= \lambda(s_i, t) \\ \lambda_t &= \exp(Y_t) + \eta Z_{t-1} + \kappa \lambda_{t-1} \\ Y_t &\sim \text{Gauss}(\alpha, (I_{n,n} - C)^{-1}M), \end{aligned} \quad (3.3)$$

که در آن $\lambda_t = (\lambda(s_1, t), \lambda(s_2, t), \dots, \lambda(s_n, t))^T$ می باشد. که λ_t یک زنجیره مارکوف در فضای $(\mathbb{R}^+)^n$ است. در ادامه تحلیل بیزی مدل بیان شده، ارائه می شود.

۱.۳ استنباط بیزی

تجزیه و تحلیل بیزی مدل *SPINGARCH* را می توان با استفاده از برنامه های نرم افزاری پیشنهادی جوزف (۲۰۱۶) انجام داد. توزیع پیشین توام پارامترهای مدل را به صورت

$$\pi(\theta) = \pi(\mu | \kappa) \pi(\kappa) \pi(\alpha) \pi(\sigma) \pi(\zeta)$$

لحاظ کرده ایم، که در آن فرض استقلال را در قسمت های مختلف به جز μ و κ به دلیل محدودیت $\mu + \kappa < 1$ ، در نظر می گیریم.

در آن صورت توزیع شرطی کامل $\theta = (\eta, \kappa, \alpha, \zeta, \sigma^2)^T$ به صورت زیر حاصل می شود:

$$\pi(\theta | Z, Y) \propto \prod_{t=1}^T \pi(Z_t | \lambda_t) \pi(\lambda_t | \lambda_{t-1}, Z_{t-1}, \theta, Y_t) \pi(Y_t | \theta) \pi(\lambda_{\cdot} | \theta) \pi(Z_{\cdot} | \lambda_{\cdot}) \pi(\theta) \quad (4.3)$$

و توزیع شرطی کامل Y به صورت

$$\pi(Y | Z, \theta) \propto \prod_{t=1}^T \pi(Z_t | \lambda_t) \pi(\lambda_t | \lambda_{t-1}, Z_{t-1}, \theta, Y_t) \pi(Y_t | \theta) \pi(\lambda_{\cdot} | \theta) \pi(Z_{\cdot} | \lambda_{\cdot}) \quad (5.3)$$

به دست می آید. برای انجام هر نوع استنباط از زنجیره مارکوف مونت کارلویی برای تولید نمونه از چگالی شرطی کامل Y استفاده می کنیم که نیاز به ارزیابی به صورت زیر دارد:

$$\begin{aligned} \log(Y | \alpha, \sigma, \zeta) &\propto \frac{-T \times n}{2} \log(2\pi) + \frac{1}{2} \log |\Sigma_f^{-1}(\theta)| \\ &\quad - \frac{1}{2} (Y - \alpha)^T \Sigma_f^{-1}(\theta) (Y - \alpha) \end{aligned} \quad (6.3)$$

توجه داشته باشید که همان طور که ساختار همسایگی فرض می شود؛ ثابت C برای همه دوره های زمانی یکسان است. بنابراین می توانیم $\Sigma_f(\theta)$ را به صورت $(I_{n \times T, n \times T} - I_{t,t} \otimes C)^{-1} I_{t,t} \otimes M$ و به عنوان ماتریس کوواریانس فضا - زمان کامل بنویسیم. کمبود ساختار کواریانس به این معناست که تنها محاسبات $\frac{1}{2} (Y - \alpha)^T \Sigma_f^{-1}(\theta) (Y - \alpha)$ است که نیاز است برای همسایگان فضایی رخ بدهد. بنابراین، دشوارترین بخش محاسبه چگالی لگاریتم، محاسبه دترمینان $|\Sigma_f^{-1}(\theta)|$

است.

با استفاده از الگوریتم‌های مونت کارلویی همانند نمونه‌گیر گیبس به همراه روش مونت کارلویی همیلتونی یا روش‌های مشابهی که در بسته نرم افزاری STAN است (جوزف، ۲۰۱۶)، این امکان را می‌دهد تا برازش مدل نسبتاً سریع انجام شود. برای انجام ارزیابی مدل به مقادیر P پیش بینی کننده گذشته تکیه می‌کنیم. این روش بعد از نمونه‌گیری پارامترها از توزیع پیشین، مجموعه داده‌های جدید را نمونه برداری می‌کند. آمارها براساس داده‌های جدید محاسبه می‌شوند و با آمارهای مجموعه داده اصلی مقایسه می‌شوند. برای هر مجموعه داده $T_1(Z)$ آمارهای فضایی موران، $T_2(Z)$ لگاریتم نسبت واریانس به میانگین، $T_3(Z)$ واریانس اختلاف $Z(s_i, t) - Z(s_i, t-1)$ و $T_4(Z)$ میانگین نمونه اتورگرسیو با یک تاخیر، را به دست می‌آوریم. T_1 و T_4 ساختار فضایی- زمانی را در داده‌ها ضبط می‌کنند؛ در حالی که T_2 و T_3 تغییرپذیری و ناهموازی در هر مکان را می‌گیرند. همان‌طور که در بخش بعد خواهیم دید هیچ یک از این آمارها برای تمایز بین $SPINGARCH$ ، $SPINGARCH$ زمان ثابت (فقط مکان) و مدل $INGARCH$ کافی نیست. در عوض، بررسی هر چهار پیشگوی گذشته معمولاً برای تمایز مناسب مدل‌ها از یکدیگر است و در مورد مقادیر بزرگتر و کوچکتر از ۰/۵ باید در برازش مدل تردید کرد.

۴ مطالعه شبیه‌سازی

هدف از شبیه‌سازی‌ها، درک تأثیر مدل آماری اشتباه است و نشان دادن چگونگی بررسی پیشگوی پسینی می‌تواند در انجام انتخاب مدل بین $SPINGARCH$ ، $SPINGARCH$ زمان ثابت و مدل $INGARCH$ کمک کند. یک شبکه 20×20 را با ساختار چهارنزدیکترین همسایگی در نقطه زمانی برای شبیه‌سازی محدوده فضایی در نظر گرفته‌ایم. ابتدا فرض می‌کنیم مدل شبیه‌سازی داده‌ها به صورت زیر است:

$$\begin{aligned} Z(s_i, t) &\sim \text{Poisson}(\lambda(s_i, t)) \\ \lambda(s_i, t) &= \exp(U(s_i)) + 0.2Z(s_i, t-1) + 0.3\lambda(s_i, t-1) \\ U(s_i)|u(N_i) &\sim N(\mu(s_i), 0.5) \\ \mu(s_i) &= 0.245 \sum_{s_j \in N_i} u(s_j) \end{aligned} \quad (1.4)$$

مدل ارائه شده در (۱.۴) فرض می‌کند که تغییر امید ریاضی در هر دوره زمانی ناشی از شمارش قبلی، امید قبلی و یک عامل متغیر فضایی است که با گذشت زمان تغییر نمی‌کند. سپس، مدل $SPINGARCH$ ، (۳.۳)، مدل $SPINGARCH$ زمان ثابت (فقط فضایی) و نیز مدل $INGARCH(1, 1)$ ، را به داده‌های تولید شده برازش می‌کنیم. بازه باورمندی بیزی برای پارامترهای هر مدل در جدول ۱ آورده شده است. از نتایج جدول ۱ در می‌یابیم که مدل واقعی، $SPINGARCH$

جدول ۱: بازه باورمندی ۹۵٪ برای برآورد پارامترها با استفاده از مدل ۱.۴ به عنوان مکانیزم تولید

$INGARCH(1, 1)$	$SPINGARCH$ زمان ثابت	$SPINGARCH$	
$(-1/8, -1/7)$	$(-1/10, 0.34)$	$(-7/9, -6/5)$	$\alpha = 0$
$(0.38, 0.40)$	$(0.17, 0.21)$	$(0.08, 0.09)$	$\eta = 0.2$
$(0.56, 0.58)$	$(0.23, 0.36)$	$(0.90, 0.92)$	$\kappa = 0.3$
	$(0.36, 0.69)$	$(0.79, 1.00)$	$\sigma^2 = 0.5$
	$(0.238, 0.249)$	$(0.248, 0.249)$	$\zeta = 0.245$

زمان ثابت، پارامترهای مدل را پوشش می‌دهد در حالی که مدل $SPINGARCH$ و $INGARCH$ مقدار κ را بیشتر نشان

می‌دهند و اصطلاحاً متورم می‌کنند. علاوه بر این، مدل $SPINGARCH$ مقادیر σ^2 و ζ را متورم می‌کنند. در ادامه از توزیع پسین برای محاسبه پیشگوی پسینی استفاده می‌کنیم، که نتایج بررسی در جدول ۲ آمده است (نیکلاس و همکاران، ۲۰۱۸). در جدول ۲ ملاحظه می‌شود که وقتی داده‌ها مطابق با مدل واقعی هستند؛ توزیع پیشین حاصل می‌تواند داده‌ها را به‌طوری جاگذاری کند که داری مشخصات مشابه با داده‌های اصلی باشند. با این حال، اگر این مدل نادرست باشد؛ این کار را نمی‌توان انجام داد. در حالی که $INGARCH$ قادر به جاگذاری هیچ یک از خصوصیات مرتبه دوم داده‌های اصلی نیست؛ $SPINGARCH$ به جز آماره‌ی $AR(1)$ برای همه به خوبی عمل می‌کند. در حالی که به نظر می‌رسد $SPINGARCH$ و $SPINGARCH$ بدون تغییر زمان قادر به تولید خواص بسیار مشابهی در مرتبه دوم هستند؛ نیکلاس و همکاران (۲۰۱۸) نشان دادند که وقتی η نزدیک لبه فضای پارامتر باشد دیگر این مورد نیست. حال یک دامنه

جدول ۲: بررسی‌های پیشگوی پسینی با شبیه‌سازی توزیع پسین

بررسی‌های پیشگوی پسینی	$SPINGARCH$	$SPINGARCH$ زمان ثابت	$INGARCH(1, 1)$
۰	۰/۳۲	۰/۳۷	آماره‌ی موران
۰	۰/۳۸	۰/۲۶	نسبت واریانس به میانگین
۰/۰۴	۰/۵۸	۰/۲	واریانس اختلاف
۱	۰/۵۸	۰/۹۰	$AR(1)$

فضایی کوچکتر، 10×10 را در نظر می‌گیریم و فرض می‌کنیم مکانیزم مولد مدل کامل $SPINGARCH$ به صورت زیر است:

$$\begin{aligned}
 Z(s_i, t) &\sim \text{Poisson}(\lambda(s_i, t)) \\
 \lambda(s_i, t) &= \exp(U(s_i)) + 0/2 Z(s_i, t-1) + 0/3 \lambda(s_i, t-1) \\
 Y(s_i) | y(N_i) &\sim N(\mu(s_i), 0/5) \\
 \mu(s_i, t) &= 0/245 \sum_{s_j \in N_i} y(s_j, t).
 \end{aligned} \tag{2.4}$$

نتایج برازش هر سه مدل در جدول ۳ و ۴ نشان داده شده است.

جدول ۳: بازه باورمندی ۹۵٪ برای برآورد پارامترها با استفاده از مدل ۱.۴

$INGARCH(1, 1)$	$SPINGARCH$ زمان ثابت	$SPINGARCH$	
(۰/۶۶, ۰/۷۶)	(۰/۶۷, ۰/۸۳)	(-۰/۰۲, ۰/۱۹)	$\alpha = 0$
(۰/۳۵, ۰/۳۹)	(۰/۳۵, ۰/۳۷)	(۰/۱۷, ۰/۲۳)	$\eta = 0/2$
(۰/۰۳, ۰/۰۸)	(۰/۰۱, ۰/۰۷)	(۰/۲۵, ۰/۳۳)	$\kappa = 0/3$
	(۰/۰۱, ۰/۰۳)	(۰/۴۳, ۰/۵۳)	$\sigma^2 = 0/5$
	(۰/۰۹, ۰/۲۴۹)	(۰/۲۴۳, ۰/۲۴۶)	$\zeta = 0/245$

همان‌طور که در جدول ۳ نشان داده شده است؛ اگر مکانیزم تولید داده واقعی شبیه به (۲.۴) باشد؛ تفاوت زیادی بین $SPINGARCH$ زمان ثابت و مدل $INGARCH$ وجود ندارد.

جدول ۴: بررسی‌های پیشگوی پسینی با شبیه‌سازی توزیع پسین

بررسی‌های پیشگوی پسینی	$SPINGARCH$	$SPINGARCH$ زمان ثابت	$INGARCH(1,1)$
آماره‌ی موران	۰/۳۸	۰	۰
نسبت واریانس به میانگین	۰/۴۳	۰	۰
واریانس اختلاف	۰/۵۶	۰	۰
$AR(1)$	۰/۵۴	۰/۳۲	۰/۳۱

بحث و نتیجه‌گیری

در این مقاله مدل‌های سری‌زمانی شمارشی بر اساس فرایند پواسن ارائه و مورد مطالعه قرار گرفتند. همچنین ساختار مدل‌های سری‌زمانی شمارشی با در نظر گرفتن همبستگی فضایی از طریق نزدیکترین همسایه بیان شد و در یک مطالعه شبیه‌سازی نحوه تولید نمونه از این مدل‌ها ارائه گردید. به‌وسیله رهیافت بیزی، پارامترهای مدل‌ها برآورد شدند و مورد تحلیل قرار گرفتند. در مطالعه شبیه‌سازی، مدل‌های $INGARCH$ ، $SPINGARCH$ و $SPINGARCH$ زمان ثابت را به‌وسیله پیشگوی پسینی مورد مقایسه قرار دادیم. نتایج بیانگر این بود که مدل $SPINGARCH$ عملکرد مناسبی نسبت به بقیه مدل‌ها داشت.

مراجع

- Brockwell, P. J., Davis, R.A. (1991), Time Series: Data Analysis and Theory, second ed, *Springer*, New York.
- Cressie, N., Wikle, C. K. (2015). Statistics for Spatio-Temporal Data. John Wiley & Sons.
- Ferland, R., Latour, A., Oraichi, D. (2006), Integer-Valued GARCH Processes, *J. Time Ser. Anal.*, **27**, 923–942.
- Heinen, A. (2003), Modelling Time Series Count Data: An Autoregressive Conditional Poisson Model, *Technical Report MPRA Paper 8113*, University Library of Munich, Germany.
- Joseph, M. (2016). Exact Sparse Car Models in Stan. Accessed: **2017-07-17**.
- Mccillagh, P., Nelder, J.A. (1989), Generalized Linear Models, second ed. Chapman & Hall, London.
- Nelder, J. A., Wedderburn, R. W. M. (1972), Generalized Linear Models, *J. R. Stat. Soc. Ser. A*, **135**, 370–384.
- Nicholas J. C., Mark, S. K., Philip M. D. (2018), A Spatially Correlated Auto-regressive Model for Count Data, stat AP.
- Priestley, M. B. (1981), Spectral Analysis and Time Series, Academic Press, London.
- Rydberg, T. H., Shephard, N. (2000), A Modeling Framework for the Prices and Times of Trades on the New York Stock Exchange. In: Fitzgerald, W.J., Smith, R.L., Walden, A.T., Young, P.C. (Eds.),

Nonlinear and Nonstationary Signal Processing. Isaac Newton Institute and Cambridge University Press, Cambridge, pp. 217–246.

Streett, S. (2000), Some Observation Driven Models for Time Series of Counts, Ph. D. thesis, Colorado State University, Department of Statistics.

Zhu, R., Joe, H. (2006), Modelling Count Data Time Series with Markov Processes Based on Binomial Thinning, *J. Time Ser. Anal*, **27**, 725–738.



برآورد تغییرنگار با استفاده از روش رگرسیون بردار پشتیبان

عباس توسلی^۱، یدا... واقعی^۱، علیرضا ناظمی^۲
^۱گروه آمار، دانشکده علوم ریاضی و آمار، دانشگاه بیرجند
^۲دانشکده علوم ریاضی، دانشگاه صنعتی شاهرود

چکیده: برازش مدل تغییرنگار یک گام اساسی در اکثر مطالعات زمین آماری است. در حال حاضر، روش های حداکثر درستنمایی و حداقل مربعات برای برازش مدل تغییرنگار مورد استفاده قرار می گیرند. روش حداکثر درستنمایی نیاز به پیش فرض هایی دارد که به سختی برقرار خواهد بود. همچنین روش حداقل مربعات نیاز به انتخاب یک مدل اولیه مناسب برای تغییرنگار دارد که در صورت انتخاب نادرست مدل اولیه، بر دقت پیشگویی نهایی اثر منفی خواهد گذاشت. اشکال دیگر روش های زمین آماری آن است که مدل های متداول و معتبر تغییرنگار رفتار هموار و افزایشی دارند؛ در حالیکه تغییرنگار داده ها اغلب نوسانات زیادی دارد و بهترین مدل هم به آن به خوبی برازش نمی شود. هدف این مطالعه استفاده از یکی از روش های یادگیری ماشین بنام «رگرسیون بردار پشتیبان» برای برآورد تغییرنگار و استفاده از آن در پیشگویی کریگیدن است. روش رگرسیون بردار پشتیبان می تواند بدون نیاز به انتخاب مدل اولیه و بصورت خودکار تغییرنگار را برآورد کند. در این مطالعه با استفاده از چندین مجموعه داده فضایی شبیه سازی شده، دقت روش رگرسیون بردار پشتیبان با روش حداقل مربعات در برآورد تغییرنگار و پیشگویی با کریگیدن مقایسه گردید. نتایج نشان داد که روش رگرسیون بردار پشتیبان می تواند رقیب مناسبی برای روش های زمین آماری در برآورد تغییرنگار باشد.

واژه های کلیدی: وابستگی فضایی، تغییرنگار، رگرسیون بردار پشتیبان، کریگیدن.

کد موضوع بندی ریاضی (۲۰۱۰): 62M45, 62J02, 62G08, 62 – 04, 62 – 07

۱ مقدمه

داده هایی که در یک محدوده فضایی (یا جغرافیایی) مشاهده شده و بر اساس موقعیتشان در آن محدوده به یکدیگر وابسته اند داده های فضایی گفته می شود. یکی از اهداف مهم در تحلیل این داده ها پیشگویی فضایی در موقعیت های جدید است. با اینحال، دستیابی به این هدف مستلزم تعیین ساختار همبستگی فضایی موجود در داده هاست. تابعی که برای بیان ساختار همبستگی داده ها مورد استفاده قرار می گیرد تابع «تغییرنگار» است که اولین بار توسط ماترن (۱۹۶۳) به این اسم نامگذاری

^۱ نام و ایمیل ارائه دهنده مقاله: عباس توسلی، Tavassoli.a@Birjand.ac.ir

شد. مدل‌بندی دقیق تغییرنگار، گامی اساسی در اکثر مطالعات زمین‌آمار است. تغییرنگار علاوه بر توصیف رفتار فضایی متغیر مورد مطالعه، تاثیر قابل توجهی در دقت پیشگویی از طریق کریگیدن دارد (دسازیز و رنارد، ۲۰۱۳). در حالت کلی، انتخاب مدل اولیه برای تغییرنگار شهودی بوده و ممکن است بهترین مدل تغییرنگار انتخاب نشود. در مورد مدل‌سازی تغییرنگار تاکنون مطالعات بسیاری صورت گرفته است. این مسئله هنوز هم در کاربردهای زمین‌آمار با دشواری‌هایی همراه است (البور و وبستر، ۲۰۱۵). در حال حاضر، دو مورد از روش‌های معمول برای مدل‌بندی تغییرنگار، روش حداقل‌درست‌نمایی و روش حداقل مربعات می‌باشند (کرسی، ۱۹۸۵، ۱۹۹۳؛ البور و وبستر، ۲۰۰۷). در روش حداقل‌درست‌نمایی، با ماکزیموم کردن تابع درست‌نمایی توام برای مقادیر مشاهده شده، مدل‌های تغییرنگار انتخاب و برازش داده می‌شود. با اینحال، روش حداقل مربعات از نظر محاسباتی ساده‌تر بوده و در اکثر بسته‌های نرم‌افزار زمین‌آمار مورد استفاده قرار گرفته است.

در تمامی روش‌هایی که تاکنون برای برازش مدل تغییرنگار بکار گرفته شده است، ساختار اولیه مدل بعنوان پیش‌فرض انتخاب و پذیرفته شده و سپس پارامترهای بهینه برای این مدل از پیش تعیین شده برآورد شده است. با اینحال، انتخاب یک مدل تغییرنگار بهینه از بین تعداد محدودی از ساختارهای اولیه که تاکنون معرفی شده‌اند مشکل است و بنابراین هیچگاه نمی‌توان یک مدل بهینه برای تغییرنگار به منظور کاهش خطای پیشگویی پیدا کرد. روش‌های مبتنی بر هوش مصنوعی می‌توانند بر اساس مقادیر مشاهده شده مدل بهینه را پیدا کنند. یکی از این روش‌ها استفاده از رگرسیون بردار پشتیبان^۱ (SVR) است. اولین بار هوانگ و همکاران (۲۰۱۲) روش SVR را برای برازش تغییرنگار مورد استفاده قرار دادند. یوکو و همکاران (۲۰۱۵) از روشی مشابه برای پیشگویی ضخامت شکاف زغال سنگ استفاده کردند. همچنین هان و همکاران (۲۰۱۶) نشان دادند که روش مبتنی بر SVR در مقایسه با روش WLS برای مدل‌بندی تغییرنگار قادر است بدون در نظر گرفتن شکل مدل اولیه، به شکل دقیق‌تری به تغییرنگار تجربی برازش یابد و دقت پیشگویی کریگیدن را افزایش دهد. با توجه به اینکه روش SVR بر اصل کمینه‌سازی ریسک ساختاری استوار است دارای توانایی بالایی در تعمیم‌پذیری بوده و فرمول‌بندی آن بر مبنای داده‌های مشاهده شده ساده است.

هدف این مطالعه بررسی عملکرد روش یادگیری ماشین SVR در برازش تغییرنگار به منظور استفاده از آن در روش کریگیدن معمولی است. بدین منظور، در بخش ۲ توضیح مختصری در رابطه با مدل‌بندی تغییرنگار ارائه شده و در بخش ۳ روش SVR معرفی شده است. در بخش ۴ با در نظر گرفتن چندین مجموعه داده فضایی شبیه‌سازی شده، به مقایسه روش SVR و روش حداقل مربعات در برازش مدل تغییرنگار پرداخته شده است. مقایسه دو روش بر اساس دقت پیشگویی کریگیدن انجام گرفته است.

۲ مدل‌بندی زمین‌آمار تغییرنگار

برای برآورد و مدل‌بندی تغییرنگار نیاز به نمادها و مفروضاتی برای داده‌های فضایی داریم. فرض کنید $\{Z(s); s \in D\}$ یک میدان تصادفی باشد بطوریکه $d \geq 1$ ، $D \subset R^d$ است. این میدان مانای ذاتی نامیده می‌شود هرگاه میانگین میدان تصادفی ثابت (مستقل از s) بوده و واریانس تفاضل بین مقادیر میدان تصادفی در دو موقعیت s و $s+h$ فقط تابعی از فاصله بین آنها باشد (محمدزاده، ۱۳۹۸). به عبارت دیگر، هرگاه برای هر $s, s+h \in D$ دو شرط زیر برقرار باشد:

$$E(Z(s) - Z(s+h)) = 0 \quad (۱)$$

$$Var(Z(s) - Z(s+h)) = 2\gamma(h) \quad (۲)$$

^۱Support Vector Regression

تابع $\gamma(h)$ که تنها تابعی از فاصله h است نیم تغییرنگار و $2\gamma(h)$ تغییرنگار نامیده می شود؛ ولی در این مقاله هرکجا واژه تغییرنگار را بکار می بریم منظور همان $\gamma(h)$ است. در روش متداول (مبتنی بر روش حداقل مربعات)، مدل سازی تغییرنگار شامل دو مرحله است: ابتدا محاسبه تغییرنگار تجربی و سپس انتخاب و برازش یک مدل تغییرنگار معتبر به تغییرنگار تجربی داده ها. **ماترن (۱۹۶۲)** تغییرنگار تجربی را، برای برآورد کردن تغییرنگار، بصورت زیر پیشنهاد کرد:

$$\hat{\gamma}(h) = \frac{1}{2|N_h|} \sum_{|N_h|} (z(s_i + h) - Z(s_i))^2 \quad (1.2)$$

بطوریکه $Z(s_i)$ و $Z(s_i + h)$ مقادیر مشاهده شده برای متغیر Z در دو موقعیت s_i و $s_i + h$ هستند. و $|N_h|$ تعداد زوج مشاهدات متمایزی است که به فاصله h از یکدیگر قرار گرفته اند. برآوردگر ماترن نااریب و سازگار است (میراندا و دمیراندا، ۲۰۱۱). در صورتی که γ تنها به طول فاصله لگ ها (و نه به جهت آنها) بستگی داشته باشد تغییرنگار همسانگرد نامیده می شود.

در دومین مرحله از برآورد تغییرنگار (یعنی انتخاب یک مدل تغییرنگار معتبر و برازش آن به تغییرنگار تجربی یا به بیان دیگر، مدل بندی تغییرنگار)، هر تابع پیوسته ای که داری خصوصیت همیشه منفی شرطی باشد می تواند بعنوان یک مدل معتبر تغییرنگار در نظر گرفته شود. این خصوصیت بیان می کند که برای هر $\{s_i \in D \subset R^d; 1 \leq i \leq m\}$ و برای هر دنباله $\{a_i \in R; 1 \leq i \leq m\}$ که $\sum_{i=1}^m a_i = 0$ ، آنگاه $\sum_{i=1}^m \sum_{j=1}^m a_i a_j \gamma(s_i - s_j) \leq 0$. این ویژگی، اعتبار تغییرنگار را تضمین می کند. پرکاربردترین توابعی که برای مدلسازی تغییرنگار بکار برده می شوند توابع نمایی، گاوسی و کروی هستند که مدل آنها در جدول ۱ نشان داده شده است (الیور و وبستر، ۲۰۰۷). هر ترکیب خطی مثبت از این مدل های تغییرنگار نیز می تواند بعنوان مدل معتبر در نظر گرفته شود.

جدول ۱: مدل های معتبر نیم تغییرنگار

مدل	نیم تغییرنگار
خطی	$\gamma(h) = c_0 + ch$
نمایی	$\gamma(h) = c_0 + c \left(1 - \exp\left(-\frac{h}{a}\right)\right), h \in R^d, d \geq 1$
گاوسی	$\gamma(h) = c_0 + c \left(1 - \exp\left(-\frac{h^2}{a^2}\right)\right), h \in R^d, d \geq 1$
کروی	$\gamma(h) = \begin{cases} c_0 + c \left(\frac{3}{2} \frac{h}{a} - \frac{1}{2} \frac{h^3}{a^3}\right) & 0 \leq h < a \\ c_0 + c & h > a \end{cases}$

در این مدل ها، c_0 اثر قطعه ای، $c_0 + c$ اِزاره و a دامنه تغییرنگار را نشان می دهند. این پارامترها از روی مشاهدات برآورد می شوند. دامنه، کوچکترین فاصله ای است که در آن مقدار $\gamma(h)$ برابر با اِزاره می شود. اگر فاصله دو مشاهده بزرگتر از مقدار دامنه باشد بر یکدیگر اثری نداشته و ناهمبسته خواهند بود. همچنین اثر قطعه ای که مقدار حدی تغییرنگار را در نقطه $h = 0$ نشان می دهد از نظر تئوری باید صفر باشد، اما در عمل بدلیل خطای نمونه گیری، خطای اندازه گیری و یا تغییرات شدید کمیت مورد بررسی در موقعیت های نزدیک به هم، ممکن است بزرگتر از صفر باشد (محمدزاده، ۱۳۹۸).

پس از انتخاب مدل اولیه برای تغییرنگار، بایستی پارامترهای مدل بهینه با کمک روش‌های مینیمم‌سازی استاندارد برآورد گردد. برای این منظور زوج نقاط $(h, \hat{y}(h))$ همانند زوج داده‌های رگرسیونی در نظر گرفته شده و پارامترهای مدل به گونه‌ای برآورد می‌شوند که اختلاف بین تغییرنگار تجربی و مدل تغییرنگار به کمترین مقدار ممکن برسد. بر اساس این ایده از یکی از روشهای حداقل مربعات خطا استفاده می‌کنیم. تاکنون چندین روش حداقل مربعات برای مدل‌بندی تغییرنگار پیشنهاد شده است؛ از جمله، روش حداقل مربعات معمولی (OLS)، روش حداقل مربعات وزنی (WLS) و روش حداقل مربعات تعمیم‌یافته (GLS) (کرسبی، ۱۹۹۳). در تمامی این روش‌ها، برازش مدل از طریق کمینه‌سازی یک تابع هزینه انجام می‌گیرد. تابع هزینه فاصله بین مدل برآورد شده و تغییرنگار تجربی را اندازه‌گیری می‌کند. روش OLS نسبت به دیگر روش‌ها ساده‌تر است؛ از طرفی، روش GLS نیاز به برآورد دقیق ماتریس کواریانس برآوردگر تغییرنگار دارد که بدست آوردن آن حتی برای فرآیند گاوسی بسیار مشکل است. مطالعات اخیر نشان داده است که نتایج بدست آمده از روش WLS در برآورد مدل تغییرنگار نسبت به دو روش دیگر رضایت‌بخش‌تر است (دسازیز و رنارد، ۲۰۱۳؛ الیور و وبستر، ۲۰۱۵) و در کنار پیشنهادهایی که برای انتخاب وزن‌ها در روش WLS شده است روش کرسبی (۱۹۸۵) بیشترین کاربرد را دارد.

حالت‌هایی وجود دارد که هیچیک از مدل‌های مذکور بخوبی به تغییرنگار تجربی برازش داده نمی‌شوند. اشکال روش‌های زمین‌آماری این است که مدل‌های متداول (و معتبر تغییرنگار) رفتار هموار و افزایشی دارند در حالیکه تغییرنگار داده‌ها اغلب نوسانات زیادی دارد و بهترین مدل هم به خوبی به آن برازش نمی‌شود. یکی از روش‌های جایگزین برای مدل‌سازی تغییرنگار، روش SVR است. این روش نیاز به انتخاب مدل اولیه برای تغییرنگار نداشته و می‌تواند مستقیماً تغییرنگار بهینه را تولید کند. در بخش بعد به معرفی این روش پرداخته شده است.

۳ روش SVR

یکی از تکنیک‌های پر کاربرد در زمینه هوش مصنوعی، ماشین بردار پشتیبان (SVM) است. SVM یکی از روشهای یادگیری با نظارت است که در ابتدا برای مسائل طبقه‌بندی طراحی شد و سپس با استفاده از تابع زیان ϵ -insensitive برای مسائل رگرسیونی توسعه داده شد. این روش در مسائل رگرسیونی به SVR شناخته می‌شود. SVR بر مبنای تئوری یادگیری آماری استوار است که از اصل کمینه‌سازی ریسک ساختاری استفاده می‌کند و در نهایت به یک جواب بهینه کلی منجر می‌شود. کمینه‌سازی ریسک ساختاری، نسبت به روش متداول کمینه‌سازی ریسک تجربی (مورد استفاده در الگوریتم‌هایی مانند شبکه‌های عصبی مصنوعی و روش‌های کلاسیک آماری)، برتری داشته و برخلاف روش‌هایی همانند شبکه عصبی مصنوعی، به جوابهای موضعی همگرا نمی‌شود.

یک مدل SVR در حالت کلی شامل متغیرهای ورودی، بردارهای پشتیبان، توابع کرنل و متغیر خروجی می‌باشد. فرض کنید $\{(x_i, y_i); i = 1, \dots, n\}$ مجموعه n نمونه باشد بطوریکه $x_i \in \mathbb{R}^p$ بردار ورودی و y_i خروجی متناظر باشد. به منظور یافتن یک ابرصفحه جداکننده، SVR فضای ورودی اولیه را به یک فضای دیگری با ابعاد بالاتر بنام فضای خصوصیت از طریق نگاشت غیرخطی ϕ انتقال می‌دهد؛ بطوریکه در فضای نگاشت یافته می‌توان از یک مدل شبیه رگرسیون خطی بصورت $f(x, w) = w \cdot \phi(x) + b$ استفاده کرد. به عبارت دیگر، متغیرهای ورودی اولیه از یک فضای با ابعاد کوچک‌تر به فضایی با ابعاد بزرگ‌تر تصویر می‌شوند تا یک مسئله رگرسیونی غیر خطی را به یک مسئله رگرسیونی خطی تبدیل کنند. در این رابطه w بردار وزن، b بایاس و $w \cdot \phi(x)$ ضرب داخلی دو بردار است. برای برآورد مقادیر w و b بایستی تابع ریسک $R = \frac{1}{n} \|w\|^2 + \frac{C}{n} \sum_{i=1}^n L_{\epsilon}(y_i, f(x_i))$ کمینه شود بطوریکه عبارت $\frac{1}{n} \|w\|^2$ میزان پیچیدگی مدل را کنترل

می‌کند و $\frac{C}{n} \sum_{i=1}^n L_{\epsilon}(y_i, f(x_i))$ ریسک تجربی است و توسط تابع زیان ϵ - insensitive یعنی

$$L_{\epsilon}(y, f(x)) = \begin{cases} 0 & |y - f(x)| < \epsilon \\ |y - f(x)| & \text{otherwise} \end{cases}$$

اندازه‌گیری می‌شود. بر اساس این تابع زیان اگر اختلاف بین مقدار پیشگویی شده و مقدار واقعی کمتر از ϵ باشد خطای پیشگویی نادیده گرفته می‌شود. برای رفع مشکل انجام عملیات در فضای با ابعاد بالا نیاز به استفاده از یک تابع کرنل است. با استفاده از تابع کرنل، مشکل چند بعدی و غیرخطی بودن نگاشت حل می‌شود. عملکرد SVR به میزان زیادی به انتخاب تابع کرنل وابسته است. از توابع کرنل شناخته شده در SVR می‌توان کرنل‌های خطی، چندجمله‌ای، گاوسی و سیگموئید را نام برد. پرکاربردترین تابع کرنل تابع گاوسی است که به آن کرنل شعاعی مبنا^۲ (RBF) نیز گفته می‌شود و بصورت $K(x, y) = e^{-\sigma \|x-y\|^2}$ است؛ بطوریکه σ ثابت مثبت بوده و در طول فرآیند آموزش تعیین می‌شود. مقادیر بهینه برای پارامترهای σ و C با استفاده از روش اعتبارسنجی متقاطع و با جستجوی شبکه‌ای در فضای پارامترها پیدا می‌شود. در مدل‌سازی تغییرنگار با روش SVR، همانند روش حداقل مربعات از زوج نقاط $(h, \hat{\gamma}(h))$ استفاده می‌شود. به عبارت دیگر، برآورد نقطه به نقطه یا لگ‌بندی که در آن تغییرنگار تجربی مستقیماً از روی مقادیر مشاهده شده و در لگ‌های مشخص محاسبه می‌شود برای هر دو روش زمین آماری و روش SVR مشترک است. تفاوت روش SVR با روش زمین آماری در این است که روش SVR برای برازش یک منحنی هموار به این مقادیر گسسته (جهت بدست آوردن تابعی پیوسته در تمامی لگ‌ها) نیازی به انتخاب مدل تغییرنگار معتبر ندارد و قادر است بر اساس مقادیر مشاهده شده، مدل بهینه را پیدا کند.

۴ مقایسه روش‌ها

در این بخش به مقایسه عملکرد روش SVR با روش حداقل مربعات در برآورد تغییرنگار پرداخته شده است. با توجه به اینکه حجم نمونه بر دقت برآورد تغییرنگار تاثیرگذار است از سه مجموعه داده فضایی شبیه‌سازی شده با حجم نمونه‌های ۴۰۰، ۹۰۰ و ۱۶۰۰ استفاده گردید. شبیه‌سازی داده‌های فضایی در محیط R به کمک تابع grf از بسته $geoR$ انجام گرفت، بطوریکه برای مدل تغییرنگار نظری از مدل کروی با مقادیر $c_0 = 0.1$ ، $c = 4$ و $a = 6$ برای پارامترها استفاده گردید. همچنین دامنه مقادیر x و y بعنوان موقعیت جغرافیایی مشاهدات، بین صفر تا ۳۰ انتخاب گردید. مقایسه روش‌ها بر اساس دقت پیشگویی کریگیدن انجام گرفت. بدین صورت که در هر یک از این سه مجموعه داده، ۳۰ درصد بصورت تصادفی بعنوان مجموعه داده تست (به منظور پیشگویی) کنار گذاشته شد و از بقیه مجموعه نقاط برای برآورد پارامترهای مدل‌ها استفاده شد. سپس با استفاده از مدل برآورد شده برای تغییرنگار، عملیات پیشگویی برای نقاط کنار گذاشته شده انجام گرفته و میزان خطای پیشگویی محاسبه گردید. به منظور استفاده از روشهای مذکور در برآورد تغییرنگار و استفاده از آنها در پیشگویی کریگیدن، مراحل زیر برای هر مجموعه داده اجرا گردید:

(۱) محاسبه تغییرنگار تجربی $\hat{\gamma}$ بر اساس رابطه (۱.۲) با استفاده از داده‌های مشاهده شده (۷۰ درصد از داده‌ها).

(۲) مدل بندی تغییرنگار با دو روش حداقل مربعات (با استفاده از مدل‌های اولیه جدول ۱) و روش SVR

(۳) محاسبه وزن‌های کریگیدن برای هر موقعیت s_i در مجموعه داده تست و سپس محاسبه مقدار پیشگویی در آن موقعیت.

(۴) محاسبه خطای پیشگویی برای هر روش (جدول ۲)

²Radial Basis Function (RBF)

برای مدل سازی تغییرنگار بر اساس روش SVR، از پکیج $e1071$ در نرم افزار R استفاده گردید بطوریکه «فاصله لگ» بعنوان ورودی و مقادیر تغییرنگار تجربی بعنوان خروجی مدل در نظر گرفته شد. همچنین از تابع کرنل شعاعی مبنا (RBF) استفاده گردید که نتایج بهتری نسبت به توابع کرنل خطی و چندجمله ای در بر داشت. انتخاب مقادیر بهینه برای پارامترهای ϵ ، C و σ بر اساس روش جستجوی شبکه ای انجام گرفت. تغییرنگار تجربی به همراه مدل های برازش یافته بر اساس دو روش مورد بررسی در شکل ۱ برای هر سه مجموعه داده شبیه سازی شده، نشان داده شده است. مشاهده می شود که برای هر سه مجموعه داده، روش SVR توانسته است بدون در نظر گرفتن مدل اولیه (در مقایسه با روش حداقل مربعات)، ارتباط غیرخطی بین فاصله لگ ها (متغیر ورودی) و تغییرنگار تجربی (متغیر خروجی) را پیدا کند.

خطای پیشگویی بر اساس دو شاخص RMSE و r^2 محاسبه گردید. بزرگ بودن مقدار r^2 و کوچک بودن مقدار RMSE برای هر روش، نشان می دهند که دقت پیشگویی بالاتر و در نتیجه عملکرد روش مورد نظر در مدل بندی تغییرنگار بهتر بوده است. این معیارها بصورت زیر هستند:

$$RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (Z_i - \hat{Z}_i)^2}, \quad r^2 = \left(\frac{\sum_{i=1}^n (Z_i - \bar{Z})(\hat{Z}_i - \bar{\hat{Z}})}{\sqrt{\sum_{i=1}^n (Z_i - \bar{Z})^2 \sum_{i=1}^n (\hat{Z}_i - \bar{\hat{Z}})^2}} \right)^2$$

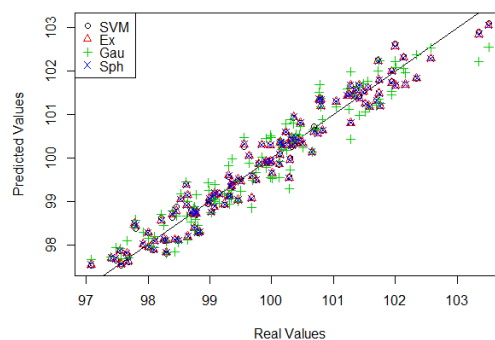
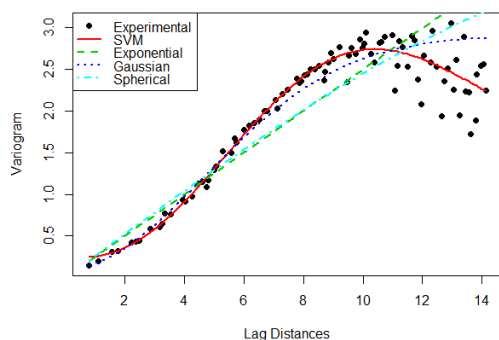
در این روابط Z_i مقادیر مشاهده شده، $\hat{Z}_i$ مقادیر پیشگویی شده، $\bar{Z}$ میانگین مقادیر مشاهده شده و $\bar{\hat{Z}}$ میانگین مقادیر پیشگویی شده است. نتایج بدست آمده برای این دو معیار، برای هر یک از روش های مورد بررسی در جدول ۲ گزارش شده است. با توجه به اینکه تولید داده های فضایی با استفاده از مدل کروی انجام گرفته است انتظار می رود که نتایج حاصل از روش حداقل مربعات بر اساس مدل اولیه کروی برای تغییرنگار، بهتر از بقیه باشد. با توجه به جدول ۲، برای حجم نمونه ۹۰۰ تایی، نتایج مدل کروی تا حدودی بهتر از دو مدل دیگر است. اما برای اولین مجموعه داده (با تعداد ۴۰۰ نمونه)، عملکرد روش SVR و برای سومین مجموعه داده، نتایج مدل نمایی بهتر بوده است.

جدول ۲: مقایسه روش های متفاوت در مدل بندی تغییرنگار بر اساس دو معیار r^2 و RMSE

$N = 1600$		$N = 900$		$N = 400$		
r^2	RMSE	r^2	RMSE	r^2	RMSE	مدل ها
۰/۷۷۶	۰/۲۴۷	۰/۸۵۶	۰/۲۸۷	۰/۹۴۶	۰/۳۲۵	SVR
۰/۷۸۵	۰/۲۴۲	۰/۸۵۹	۰/۲۸۴	۰/۹۴۵	۰/۳۲۷	مدل نمایی
۰/۷۵۳	۰/۲۵۹	۰/۸۱۰	۰/۳۳۱	۰/۹۱۰	۰/۴۲۰	مدل گاوسی
۰/۷۸۱	۰/۲۴۵	۰/۸۶۱	۰/۲۸۲	۰/۹۴۵	۰/۳۲۸	مدل کروی

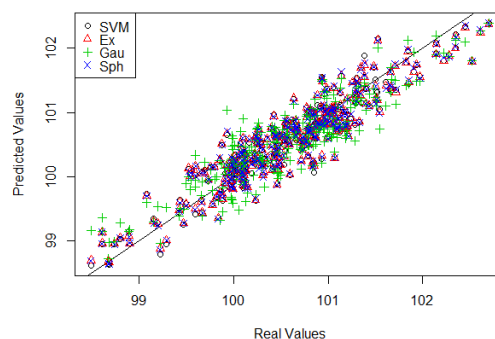
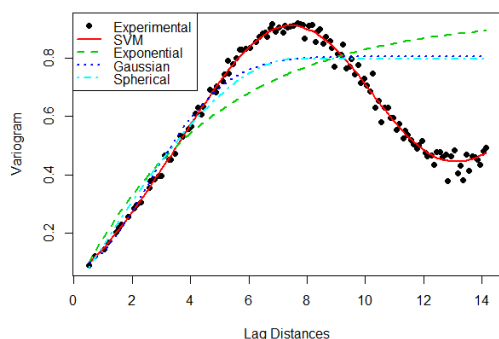
بحث و نتیجه گیری

در این مطالعه، به کمک مجموعه داده های فضایی شبیه سازی شده، عملکرد روش SVR در برآورد تغییرنگار با روش حداقل مربعات بر اساس سه مدل اولیه نمایی، گاوسی و کروی مقایسه گردید. نتایج نشان داد که روش SVR قادر است بدون در نظر گرفتن مدل اولیه، به تغییرنگار تجربی بصورت دقیق برازش یافته و بر دقت نتایج پیشگویی کریگیدن تاثیرگذار باشد. بنابراین روش SVR می تواند رقیب مناسبی برای روش های کلاسیک در برآورد تغییرنگار باشد. در این مطالعه فقط از سه مجموعه داده شبیه سازی شده با حجم نمونه های متفاوت استفاده گردید. بررسی حالت های دیگر از جمله وجود یا عدم وجود



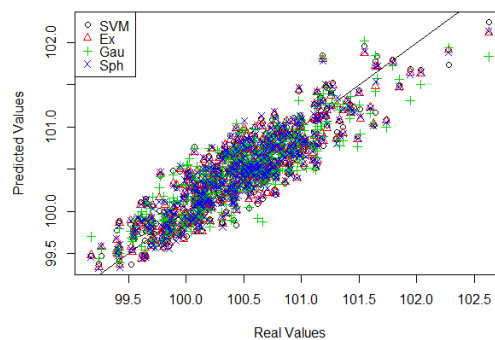
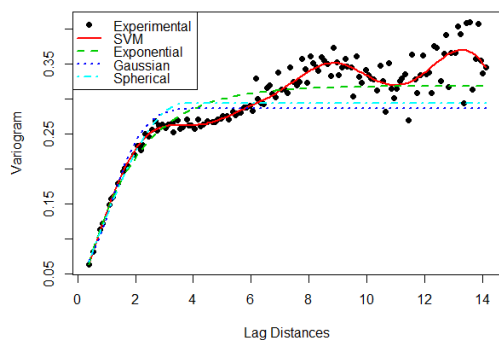
ب- تغییرنگار تجربی به همراه مدل‌های برازش شده

الف- نمودار پراکنش مقادیر واقعی در برابر پیشگویی ($N = 400$)



د- تغییرنگار تجربی به همراه مدل‌های برازش شده

ج- نمودار پراکنش مقادیر واقعی در برابر پیشگویی ($N = 900$)



و- تغییرنگار تجربی به همراه مدل‌های برازش شده

ه- نمودار پراکنش مقادیر واقعی در برابر پیشگویی ($N = 1600$)

شکل ۱: نمودار تراکنش مقادیر واقعی در برابر پیشگویی (راست)، تغییرنگار تجربی به همراه مدل‌های برازش شده (چپ)

اثر قطعه‌ای، میزان وابستگی فضایی میان داده‌ها و یا همسانگردی یا ناهمسانگردی میدان تصادفی می‌تواند در مطالعات آینده مورد بررسی قرار گیرد.

مراجع

- محمدزاده، م.، (۱۳۹۸)، آمار فضایی و کاربردهای آن، چاپ سوم، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران.
- Cressie, N. (1985), Fitting variogram models by weighted least squares. *Journal of the International Association for Mathematical Geology*, 17(5), 563-586.
- Cressie, N. (1993), *Statistics for Spatial Data*, John Wiley, New York.
- Desassis, N., Renard, D. (2013), Automatic variogram modeling by iterative least squares: univariate and multivariate cases. *Mathematical geosciences*, 45(4), 453-470.
- Han, C., Wang, J., Zheng, M., Wang, E., Xia, J., Li, G., Choe, S. (2016). New variogram modeling method using MGGP and SVR. *Earth Science Informatics*, 9(2), 197-213.
- Huang, Z., Wang, H., Zhang, R. (2012), An improved kriging interpolation technique based on SVM and its recovery experiment in oceanic missing data. *American Journal of Computational Mathematics*, 2(01), 56.
- Matheron, G. (1962), *Traité de géostatistique appliquée*. (Vol. 1). Editions Technip.
- Matheron, G. (1963), Principles of geostatistics. *Economic geology*, 58(8), 1246-1266.
- Miranda, H., de Miranda, M. S. (2011). Combining Robustness with Efficiency in the Estimation of the Variogram. *Mathematical Geosciences*, 43(3), 363-377.
- Oliver, M. A., Webster, R. (2015), *Basic steps in geostatistics: the variogram and kriging* (Vol. 106). New York: Springer.
- Webster, R., Oliver, M. A. (2007) *Geostatistics for environmental scientists*. John Wiley and Sons.
- Youkuo, C., Yongguo, Y., Wangwen, W. (2015). Coal seam thickness prediction based on least squares support vector machines and kriging method. *Electron. J. Geotech. Eng*, 20, 167-176.



تحلیل بیزی تقریبی مدل‌های رگرسیونی فضایی

فاطمه حسینی^۱، امید کریمی
گروه آمار، دانشگاه سمنان

چکیده: در مطالعات اقتصادی اغلب داده‌ها دارای وابستگی فضایی هستند و برای مدل‌بندی این نوع از داده‌ها معمولاً از سه مدل تاخیر فضایی، دوربین فضایی و خطای فضایی به عنوان سه مدل معروف رگرسیون فضایی استفاده می‌شود. هدف اصلی در مطالعه این مدل‌ها به دست آوردن برآورد پارامترها و سپس پیشگویی در موقعیت‌های جدید است. برای این منظور، در این مقاله ابتدا رهیافت ماکسیمم درستنمایی بررسی شده است و در راستای بالا بردن دقت برآورد پارامترها و کاهش زمان محاسبات، رهیافت بیزی تقریبی برای این سه مدل رگرسیونی فضایی بیان می‌شود. در نهایت در یک مثال واقعی به مقایسه عملکرد سه مدل و دو رهیافت پرداخته می‌شود.

واژه‌های کلیدی: مدل تاخیر فضایی، مدل دوربین فضایی، مدل خطای فضایی، رهیافت بیزی تقریبی.
کد موضوع‌بندی ریاضی (۲۰۱۰): 91B72, 62F15, 91G70

۱ مقدمه

معمولاً برای تحلیل داده‌های اقتصادی که وابستگی مکانی دارند از مدل‌های رگرسیون فضایی استفاده می‌شود. سه مدل معروف رگرسیون فضایی که اغلب برای این نوع داده‌ها در اقتصادسنجی فضایی مورد استفاده قرار می‌گیرد. مدل‌های اتورگرسیون فضایی یا مدل تاخیر فضایی^۱، (SLM) مدل دوربین فضایی^۲ (SDM) و مدل خطای فضایی^۳ (SEM) هستند. در مدل خطای فضایی همبستگی فضایی در جمله خطا منظور می‌شود، در مدل اتورگرسیون فضایی اثر فضایی در متغیر پاسخ منظور می‌شود و در مدل دوربین فضایی هم از طریق متغیر پاسخ و هم از طریق متغیرهای کمکی در مدل وارد می‌شود. اولین بار **انسلین (۱۹۸۸)** به معرفی اقتصادسنجی فضایی پرداخت و **داویدسون و مکینون (۱۹۹۳)**، **لسج**

^۱ Spatial Autoregressive Model or Spatial Lag Model

^۲ Spatial Durbin Model

^۳ Spatial Error Model

(۱۹۹۹)، لسج و پس (۲۰۰۹) و لسج و فیشر (۲۰۰۸) به طور مفصل تر روش های اقتصادسنجی فضایی را مورد مطالعه قرار دادند. انسلین (۲۰۱۰) برای داده ها در طول زمان طولانی و با به کار بردن مدل های رگرسیونی فضایی به بررسی و مطالعه اقتصادسنجی فضایی پرداخت. بایوند و همکاران (۲۰۱۴) به معرفی تکنیک های تحلیل داده های فضایی و کاربرد این تکنیک ها در مطالعات اقتصادسنجی با نرم افزار R پرداختند. اربیا (۲۰۱۶) و کلیجیان و پیراس (۲۰۱۷) انواع روش های اقتصادسنجی را به اقتصادسنجی فضایی تعمیم داده اند. برای تحلیل این مدل ها در اکثر مطالعات اشاره شده معمولاً انواع روش های برآورد ماکسیمم درستنمایی پیشنهاد شده است. به دلیل اضافه کردن همبستگی فضایی به این مدل ها و اضافه شدن تعداد پارامترهای مدل برای استفاده از رهیافت ماکسیمم درستنمایی تعدادی از پارامترها معلوم فرض می شوند و سایر پارامترها را برآورد و مجدداً با استفاده از برآوردهای به دست آمده پارامترهایی که معلوم فرض شدند را برآورد می کنند. این یکی از مشکلات اساسی در استفاده از رهیافت ماکسیمم درستنمایی در این مدل ها است. در رهیافت بیزی برای تمام پارامترها توزیع پیشین فرض می شود و با استفاده از تابع درستنمایی تشکیل شده و توزیع های پیشین فرض شده، توزیع پسین تشکیل و با استفاده از الگوریتم های موجود مثل الگوریتم مونت کارلوی زنجیر مارکوفی به طور همزمان پارامترها برآورد می شوند. از طرفی اجرای الگوریتم های مونت کارلویی زمان بر است و گاهی رسیدن به همگرایی در این الگوریتم ها نیاز به تکرار زیاد دارد و به کندی انجام می شود. در این مقاله از یک رهیافت بیزی تقریبی برای تحلیل مدل ها استفاده می شود و در یک مثال روش ها مورد مقایسه قرار می گیرند. ساختار مقاله به این صورت است که در بخش دوم سه مدل رگرسیون فضایی معرفی می شوند. در بخش سه تحلیل درستنمایی و در بخش چهار تحلیل بیزی تقریبی سه مدل بیان و در بخش چهارم و پنجم برای مقایسه بین دو رهیافت درستنمایی و بیزی تقریبی به تحلیل یک مجموعه داده واقعی مربوط به قیمت مسکن در شهر تهران پرداخته می شود. در نهایت بحث و نتیجه گیری ارائه شده است.

۲ مدل های رگرسیونی فضایی

فرض کنید $y' = (y_1, \dots, y_n)$ بردار مشاهدات در n موقعیت فضایی $s_1, \dots, s_n$ است و X ماتریس طرح با p متغیر کمکی و $\beta = (\beta_1, \dots, \beta_p)'$ پارامترهای رگرسیونی هستند، مدل تاخیر فضایی که در لسج و پس (۲۰۰۹) ارائه شده است، متغیر پاسخ به مجموع وزنی پاسخ ها در همسایگی هایش، متغیرهای کمکی و جمله خطا بستگی دارد و به صورت

$$y = \rho W y + X\beta + e', \quad e' \sim MVN(0, \sigma^2 I_n), \quad (1.2)$$

است، که در آن I_n یک ماتریس واحد از بعد $n \times n$ و W یک ماتریس مجاورت سطری استاندارد شده و ρ پارامتر خودهمبستگی فضایی است. مدل خطای فضایی به صورت

$$\begin{aligned} y &= X\beta + u, \quad u \sim MVN(0, \sigma^2 (I_n - \lambda W)^{-1} (I_n - \lambda W)') \\ u &= \lambda W u + e, \quad e \sim MVN(0, \sigma^2 I_n), \\ y &= X\beta + (I_n - \lambda W)^{-1} e, \end{aligned} \quad (2.2)$$

بیان می شود، که در آن λ پارامتر خودهمبستگی فضایی است. مدل دوربین فضایی به صورت

$$y = \rho W y + X\beta + WX\gamma + e', \quad e' \sim MVN(0, \sigma^2 I_n), \quad (3.2)$$

تعریف می شود، که در آن WX متغیرهای کمکی تاخیر فضایی و γ بردار ضرایب این متغیرها هستند. در این مدل پاسخ علاوه بر این که به پاسخ ها در همسایگی اش و متغیرهای کمکی بستگی دارد به مقادیر متغیرهای کمکی همسایگی اش وابسته

است. در تمام مدل‌های معرفی شده، جمله خطا گاوسی با میانگین صفر و ماتریس واریانس کوواریانس به شکل اتو رگرسیو فضایی^۴ (SAR) در نظر گرفته شده است. مدل‌های SLM و SDM ساختار پیچیده‌تری در قسمت خطی مدل که مربوط به متغیرهای کمکی می‌باشد، دارند. همچنین پارامترهای خودهمبستگی فضایی ρ در بازه $(\frac{1}{\lambda_{min}}, \frac{1}{\lambda_{max}})$ محدود می‌شود که در آن λ_{max} و λ_{min} ، کوچک‌ترین و بزرگ‌ترین مقدار ویژه W هستند. در این سه مدل ارتباط فضایی متغیرها به صورت دو به دو و به فرم عددی توسط ماتریس W مشخص می‌شود که معمولاً و نه لزوماً همیشه ماتریسی متقارن در نظر گرفته می‌شود و در آن عنصر W_{ij} ارتباط فضایی مربوط به متغیر پاسخ در موقعیت i با متغیر پاسخ در موقعیت j است. به عنوان مثال در برخی مطالعات مربوط به قیمت مسکن فاصله مرکز منطقه i از منطقه j با d_{ij} نشان می‌دهند و W_{ij} را از رابطه $\frac{1}{d_{ij}}$ به دست می‌آورند به طوری که با افزایش فاصله دو منطقه اثر فضایی کم می‌شود. **لسج و پس (۲۰۰۹)** نشان دادند که می‌توان W را به صورت یک ماتریس سطری تصادفی ساخت، به طوری که بردار تاخیر فضایی Wy شامل مقادیر ساخته شده از متوسط مشاهدات همسایه می‌باشد. منظور از ماتریس سطری تصادفی، ماتریس نامنفی است که عناصر روی هر سطر آن نرمال‌سازی شده و بنابراین مجموع هر سطر یک می‌باشد و در این صورت $\lambda_{max} = 1$ و $\rho < 1$ است.

۳ تحلیل درست‌نمایی مدل‌ها

برای به دست آوردن برآورد ماکسیمم درست‌نمایی پارامترهای مدل تاخیر فضایی، مدل (۱.۲) را می‌توان با فرض معلوم بودن ρ به صورت

$$y - \rho Wy = X\beta + \epsilon, \quad \epsilon \sim N(0, \sigma^2 I_n),$$

نوشت. بنابراین $\epsilon = (I - \rho W)y - X\beta$ و داریم

$$L(\sigma^2, \epsilon) = \left(\frac{1}{\sqrt{\pi}\sigma^2}\right)^{\frac{n}{2}} \exp\left\{-\frac{1}{\sqrt{\pi}\sigma^2} \epsilon' \epsilon\right\},$$

پس تابع درست‌نمایی مشاهدات به صورت

$$L(\rho, \beta, \sigma^2 | y) = \left(\frac{1}{\sqrt{\pi}\sigma^2}\right)^{\frac{n}{2}} |J| \times \exp\left\{-\frac{1}{\sqrt{\pi}\sigma^2} [(I - \rho W)y - X\beta]' [(I - \rho W)y - X\beta]\right\},$$

است، که در آن $|J| = \left|\frac{\partial \epsilon}{\partial y}\right| = |I_n - \rho W|$. با مشتق‌گیری از لگاریتم تابع درست‌نمایی نسبت به β و σ^2 برآورد ماکسیمم درست‌نمایی آن‌ها به صورت

$$\begin{aligned} \hat{\beta} &= (X'X)^{-1} X'(I_n - \rho W)y, \\ \hat{\sigma}^2 &= \frac{1}{n} [(I_n - \rho W)y - X\hat{\beta}]' [(I_n - \rho W)y - X\hat{\beta}] \\ &= \frac{1}{n} (y - X\hat{\beta})' (I_n - \rho W)' (I_n - \rho W) (y - X\hat{\beta}), \end{aligned}$$

به دست می‌آیند. با جایگذاری $\hat{\beta}$ در $\hat{\sigma}^2$ و با تعریف

$$A = [y - X(X'X)^{-1} X'y - \rho(Wy - (X'X)^{-1} X'Wy)],$$

داریم $\hat{\sigma}^2 = \frac{1}{n} A'A$ با قرار دادن $e_0 = y - X(X'X)^{-1} X'y$ یعنی باقیمانده رگرسیون معمولی y روی X و $e = Wy - (X'X)^{-1} X'Wy$ باقیمانده رگرسیون معمولی Wy روی X داریم

$$\hat{\sigma}^2 = \frac{1}{n} (e_0 - \rho e)' (e_0 - \rho e). \quad (1.3)$$

⁴Spatial Autoregressive Model

با جایگذاری $\hat{\beta}$ و رابطه (۱.۳) در تابع درستنمایی که اکنون تابع درستنمایی متمرکز شده^۵ نامیده می شود به صورت

$$\ell_c(\rho) = -\frac{n}{2} \ln 2\pi - \frac{n}{2} \ln \frac{1}{n} (e_0 - \rho e_1)' (e_0 - \rho e_1) + \sum_i \ln(1 - \rho \omega_i), \quad (2.3)$$

خواهد بود. پس

$$\frac{\partial \ell_c(\rho)}{\partial \rho} = \frac{n(e_0' e_1 - \rho e_1' e_1)}{e_0' e_0 - 2\rho e_0' e_1 + \rho^2 e_1' e_1} - \frac{\sum_i \omega_i}{1 - \rho \omega_i} = 0,$$

به طوری که برآورد ρ باید در بازه $(\lambda_{min}^{-1}, \lambda_{max}^{-1})$ باشد. برای این منظور **باری و پس** (۱۹۹۷) پیشنهاد کردند که یک بردار $1 \times q$ از مقادیری که در بازه $[\rho_{min}, \rho_{max}]$ است برای ρ در نظر گرفته شود و آن گاه

$$\begin{pmatrix} \ell(\rho_1) \\ \vdots \\ \ell(\rho_q) \end{pmatrix} \propto \begin{pmatrix} \sum_i \ln(1 - \rho_1 \omega_i) \\ \vdots \\ \sum_i \ln(1 - \rho_q \omega_i) \end{pmatrix} - \frac{n}{2} \begin{pmatrix} \ln e'(\rho_1) e(\rho_1) \\ \vdots \\ \ln e'(\rho_q) e(\rho_q) \end{pmatrix},$$

که در آن $e(\rho) = e_0 - \rho e_1$ و $\rho_1 \dots \rho_q$ مقادیری هستند که بهینه سازی حول این مقادیر انجام می شود. با به کار بردن یک الگوریتم تکرار شونده و با در نظر گرفتن دقت مورد نظر یک شبکه به اندازه ی کافی خوب از q مقدار برای ρ ، $l(\rho)$ مشخص و مقدار بهینه تعیین خواهد شد. پس از تعیین برآورد ماکسیمم درستنمایی ρ به صورت $\hat{\rho}$ آن گاه می توان $\hat{\beta}$ و $\hat{\sigma}^2$ را به دست آورد. برآورد درستنمایی پارامترهای دو مدل خطای فضایی و دوربین فضایی را می توان به طور مشابه به دست آورد.

۴ تحلیل بیزی معمولی و تقریبی مدل های رگرسیون فضایی

فرض کنید مدل تاخیر فضایی به صورت

$$y = \rho_{Lag} W y + X \beta + e, \quad (1.4)$$

است، به طوری که $e \sim MVN(0, \sigma^2 I_n)$. برای β یک توزیع پیشین ناآگاهی بخش از نوع مبهم، یعنی یک توزیع نرمال با واریانس بزرگ و برای σ^2 توزیع پیشین گامای وارون فرض می شوند. بنابراین

$$\begin{aligned} \pi(\beta, \sigma^2) &= \pi(\beta | \sigma^2) \pi(\sigma^2) \\ &= N(c, \sigma^2 T) IG(a, b) \\ &= \frac{b^a}{(\sqrt{2\pi})^{\frac{k}{2}} |T|^{\frac{1}{2}} \Gamma(a)} (\sigma^2)^{-(a + \frac{k}{2} + 1)} \\ &\times \exp\left[-\frac{\{B - c\}' T^{-1} (B - c) + 2b}{2\sigma^2}\right], \\ \pi(\sigma^2) &= \frac{b^a}{\Gamma(a)} (\sigma^2)^{-(a+1)} \exp\left(-\frac{b}{\sigma^2}\right), \quad \sigma^2 > 0, a, b > 0 \end{aligned}$$

و برای پارامتر ρ_{Lag} توزیع یکنواخت به صورت $\pi(\rho_{Lag}) \sim U(\lambda_{min}^{-1}, \lambda_{max}^{-1})$ که در آن λ_{min} و λ_{max} کوچک ترین و بزرگ ترین مقدار ویژه W هستند.

فرض کنید تمام مشاهدات به صورت $D = \{y, X, W\}$ و پارامترهای مدل به صورت $\theta = \{\beta, \sigma^2, \rho_{Lag}\}$ در نظر گرفته شوند. تابع درستنمایی برای این مدل به صورت $P(D | \beta, \sigma^2, \rho) = (2\pi\sigma^2)^{-\frac{n}{2}} |A| \exp\{-\frac{1}{2\sigma^2} (Ay - X\beta)' (Ay -$

⁵Function Likelihood Concentrated

$\{X|\beta\}$ است، که در آن $A = (I_n - \rho W)$. بنابراین توزیع پسین به صورت

$$\pi(\beta, \sigma^2, \rho_{Lag}|D) \propto \pi(D|\beta, \sigma^2, \rho_{Lag})\pi(\beta|\sigma^2)\pi(\sigma^2)\pi(\rho_{Lag}), \quad (2.4)$$

است، که شکل پیچیده‌ای دارد و برای به دست آوردن برآوردهای بیزی پارامترها از الگوریتم‌های $MCMC$ استفاده می‌شود که اجرای الگوریتم‌ها بسیار زمان‌بر و گاهی اوقات برای رسیدن به همگرایی نیاز به صرف زمان بسیار است. به عنوان یک پیشنهاد می‌توان از رهیافت بیزی تقریبی استفاده کرد. [رو و همکاران \(۲۰۰۹\)](#)، [رو و هلد \(۲۰۰۵\)](#) و [ایدسویک و همکاران \(۲۰۱۲\)](#) با فرض نرمال بودن متغیرهای پنهان فضایی در مدل‌های شامل متغیرهای پنهان گاوسی روش بیز تقریبی را به جای روش بیز معمولی و الگوریتم‌های نمونه‌گیری مونت کارلویی معرفی کرده و نشان دادند محاسبات این روش سریع‌تر و نتایجی معادل روش‌های بیزی معمولی ارائه می‌کند. [ایدسویک و همکاران \(۲۰۱۲\)](#) با در نظر گرفتن توزیع پیشین نرمال برای متغیرهای پنهان در مدل‌های شامل متغیر پنهان طبق قضیه زیر نشان دادند که توزیع پسین تقریبی این متغیرها متعلق به خانواده توزیع نرمال است.

قضیه ۱.۴. اگر در مدل، متغیرهای پنهان x دارای توزیع پیشین نرمال $x \sim N(H\beta, \Sigma_\theta)$ و متغیر پاسخ y متعلق به خانواده نمایی $\pi(y_i|x_i) = \exp\{y_i x_i - b(x_i) + c(y_i)\}$ باشد، با خطی کردن حول یک مقدار ثابت، توزیع را می‌توان با توزیع نرمال به صورت

$$\hat{\pi}(x|y, \eta) \approx N_n(\hat{\mu}_{x|y}(\mathbf{y}, \mathbf{x}^\circ), \hat{\Sigma}_{x|y, \eta}(\mathbf{x}^\circ)),$$

تقریب زد که در آن

$$\begin{aligned} \hat{\mu}_{x|y}(\mathbf{y}, \mathbf{x}^\circ) &= H\beta + \Sigma_\theta A' R^{-1}(z(\mathbf{y}, \mathbf{x}^{obs}) - AH\beta), \\ \hat{\Sigma}_{x|y, \eta}(\mathbf{x}^\circ) &= \Sigma_\theta - \Sigma_\theta' R^{-1} A \Sigma_\theta, \quad R = A \Sigma_\theta A' + P, \end{aligned}$$

و P یک ماتریس قطری با عناصر $P(i, i) = \frac{1}{b''(x_i)} i = 1, \dots, k$ به ازای $z_i(y_i, x_i^\circ) = [y_i - b'(x_i^\circ) + x_i b''(x_i^\circ)] / b(x_i^\circ)$ می‌باشند.

برهان: مراجعه شود به [ایدسویک و همکاران \(۲۰۱۲\)](#).

برای تعیین مقدار ثابت و خطی‌سازی قسمت درستنمایی در آن و برازش یک تقریب نرمال به پسین $\hat{\pi}(x|y, \eta)$ از الگوریتم زیر استفاده می‌شود.

الف- مقدار اولیه x^* به عنوان مثال $x^* = M(x|\eta)$ ، انتخاب و $m = 0$ قرار داده شود.

ب- از قضیه ۱.۴ توزیع پسین تقریبی در x^m ، یعنی $\hat{\pi}(x|y, \eta) \equiv N_n(\hat{\mu}_{x|y, \eta}(x^{(m)}), \hat{\Sigma}_{x|y, \eta}(x^{obs}))$ تعیین می‌شود.

ج- قرار دهید $x^{(m+1)} = \hat{M}(x|y, \eta)$

د- $m = m + 1$ اختیار شود و تا رسیدن به همگرایی، الگوریتم از مرحله (ب) مجدداً شروع شود. پس از رسیدن به همگرایی، الگوریتم متوقف و توزیع پسین تقریبی تعیین می‌شود.

برای برآورد بیزی پارامترها توزیع پسین پارامترهای مدل به صورت

$$\pi(\eta|y) = \frac{\pi(y|x)\pi(x|\eta)\pi(\eta)}{\pi(y)\pi(x|y, \eta)} \propto \frac{\pi(y|x)\pi(x|\eta)\pi(\eta)}{\pi(x|y, \eta)},$$

است، که برای هر مقدار x ، قابل محاسبه می‌باشد. برای مثال می‌توان از مد توزیع پسین تقریبی $\hat{\pi}(x|y, \eta)$ برای یک بردار ثابت از پارامترها، یعنی $x = \hat{M}(x|y, \eta)$ استفاده کرد. با به کار بردن تقریب نرمال برازش شده به توزیع شرطی کامل $\pi(x|y, \eta)$ ، یک چگالی تقریبی برای توزیع پسین پارامترهای مدل به صورت

$$\hat{\pi}(\eta|y) \propto \frac{\pi(y|x)\pi(x|\eta)\pi(\eta)}{\hat{\pi}(x|y, \eta)} \Big|_{x=\hat{M}(x|y, \eta)}, \quad (3.4)$$

حاصل می‌شود، که این تقریب توسط **رو و همکاران (۲۰۰۹)** یک تقریب لاپلاس از چگالی پسین پارامترهای مدل نامیده شد. برای تعیین پسین تقریبی (۳.۴) از الگوریتم زیر استفاده می‌شود:

الف- مد تابع $\log \hat{\pi}(\eta|y)$ با روش بهینه‌سازی شبه-نیوتن یا نیوتن رافسون تعیین و با η^* نشان داده می‌شود.

ب- ماتریس هسیان از رابطه $H = \left(\frac{\partial^2 \log \hat{\pi}(\eta|y)}{\partial \eta^i \partial \eta^j} \right)$ ^۶ محاسبه و معکوس آن به صورت $H^{-1} = V\Lambda V'$ تجزیه شود، که در آن V ماتریس بردارهای ویژه و Λ ماتریس قطری مقادیر ویژه آن است.

ج- مبدا مختصات به مد η^* انتقال داده شود، فرمول مختصات در مبدا η^* به صورت $\eta(t) = \eta^* + V\Lambda^{-1}t$ تعریف می‌شود، که در آن t مقادیر استاندارد شده هستند. به عنوان مثال، برای حالت دو بعدی $t = (t_1, t_2)$ ، با قرار دادن مبدا مختصات $t = (0, 0)$ در رابطه بالا $\eta(0) = \eta^*$ حاصل می‌شود.

د- با شروع از مبدا مختصات جدید روی هر یک از محورها نقاطی به فاصله مقادیر صحیح δ_t به گونه‌ای اختیار شوند، که شرط

$$\log \hat{\pi}(\eta(0)|y) - \log \hat{\pi}(\eta(t)|y) < \delta_t, \quad (4.4)$$

برقرار باشد، سپس به طور مشابه نقاط درون صفحات نیز تعیین شوند. معمولاً در نامساوی بالا $\delta_t = 2/5$ در نظر گرفته می‌شود. با به کار بردن این الگوریتم η_i هایی از توزیع $\pi(\eta|y)$ تولید می‌شوند که می‌توان از آن‌ها برای تقریب توزیع $\pi(\eta|y)$ استفاده کرد.

برای تعمیم رهیافت بیزی تقریبی به مدل‌های رگرسیونی فضایی در مطالعات اقتصادسنجی فضایی ابتدا مدلی به صورت یک مدل آمیخته خطی تعمیم یافته شامل متغیرهای پنهان فضایی مدلی بندی می‌شود. در راستای کاهش زمان محاسبات متغیرهای پنهان موجود در مدل با یک میدان تصادفی مارکوفی گاوسی مدلی بندی می‌شوند. ماتریس واریانس-کوواریانس میدان تصادفی مارکوفی گاوسی به صورت یک ماتریس دقت^۷ که عکس ماتریسی واریانس-کوواریانس میدان تصادفی گاوسی است بیان می‌شود که یک ماتریس پراکنده^۸ است و همین باعث سادگی در محاسبات و کاهش زمان محاسبات می‌شود. فرض کنید $y' = (y_1, \dots, y_n)$ بردار متغیر پاسخ در n موقعیت $s_1, \dots, s_n$ باشد، به طوری که y می‌تواند گاوسی یا غیرگاوسی باشد و متعلق به خانواده نمایی با میانگین μ_i در نظر گرفته می‌شود. اکنون می‌توان سه مدل رگرسیون فضایی تاخیر فضایی، خطای فضایی و دوربین فضایی را به صورت مدل کلی

$$\eta_i = g(\mu_i) = z_i'\beta + u_i, \quad (5.4)$$

در نظر گرفت، که در آن پیشگویی کننده η_i شامل اثرات ثابت z_i و تصادفی u_i است، به طوری که این مدل می‌تواند شامل اثرات غیرخطی از متغیرهای پنهان و اثرات غیرخطی از متغیرهای کمکی نیز باشد. فرض کنید بردار θ_1 ابر پارامترهای

^۶Hessian Matrix

^۷Precision Matrix

^۸Sparse Matrix

مربوط به y و θ_2 بردار ابر پارامترهای مربوط به بردار x است، به طوری که بردار x شامل اثر کلی پیشگویی کننده خطی برای هر مشاهده، پارامترهای رگرسیونی و اثرات پنهان فضایی به صورت $x' = (\eta_i, \beta', u')$ باشد. بردار x به صورت یک میدان تصادفی مارکوفی گاوسی (GMRF) با ماتریس دقت $Q(\theta_2)$ فرض می شود. بنابراین تابع درستنمایی مدل با در نظر گرفتن $\theta = (\theta_1, \theta_2)$ به شکل $\pi(y|x, \theta) = \prod_{i \in I} \pi(y_i|x_i, \theta)$ که در آن $1 \leq l \leq n$ ، بیانگر مشاهدات به غیر از گمشده‌ها است. توزیع توام x یعنی پارامترهای مدل، اثرات پنهان و ابر پارامترها یعنی بردار θ به صورت

$$\begin{aligned} \pi(x, \theta|y) &\propto \pi(\theta)\pi(x|\theta) \prod_{i \in I} \pi(y_i|x_i, \theta) \\ &\propto \pi(\theta) |Q(\theta)|^{\frac{n}{2}} \exp \left\{ -\frac{1}{2} x' Q(\theta) x + \sum_{i \in I} \log(\pi(y_i|x_i, \theta)) \right\}, \end{aligned} \quad (6.4)$$

توزیع‌های حاشیه‌ای برای اثرات پنهان $\pi(x_i|y) \propto \int \pi(x_i|\theta) \pi(\theta|y) d\theta$ و برای ابر پارامتر $\pi(\theta_j|y) \propto \int \pi(\theta|y) d\theta_{-j}$ به دست می آیند، که در آن θ_{-j} بردار ابر پارامترها بدون θ_j است. اکنون از رهیافت بیزی تقریبی دیدیم $\pi(x|y, \theta)$ را می توان با توزیع نرمال $\hat{\pi}(x|y, \theta)$ تقریب زد و از خواص توزیع نرمال $\pi(x_i|\theta, y)$ به دست آورد. $\pi(\theta|y)$ را می توان از الگوریتم توضیح داده شده در بالا تعیین نمود. بنابراین تقریبی از توزیع حاشیه‌ای برای x به صورت

$$\hat{\pi}(x_i|y) = \sum_k \hat{\pi}(x_i|\theta_k, y) \hat{\pi}(\theta_k|y) \Delta_k,$$

محاسبه می شود، که در آن k تعداد نمونه‌های تولید شده از $\pi(\theta|y)$ از الگوریتم توضیح داده شده در بالا تعیین می شود و Δ_k وزن تخصیص داده شده به هر یک از عناصر بردار θ_k است که معمولاً یکسان و برابر ۱ در نظر گرفته می شود.

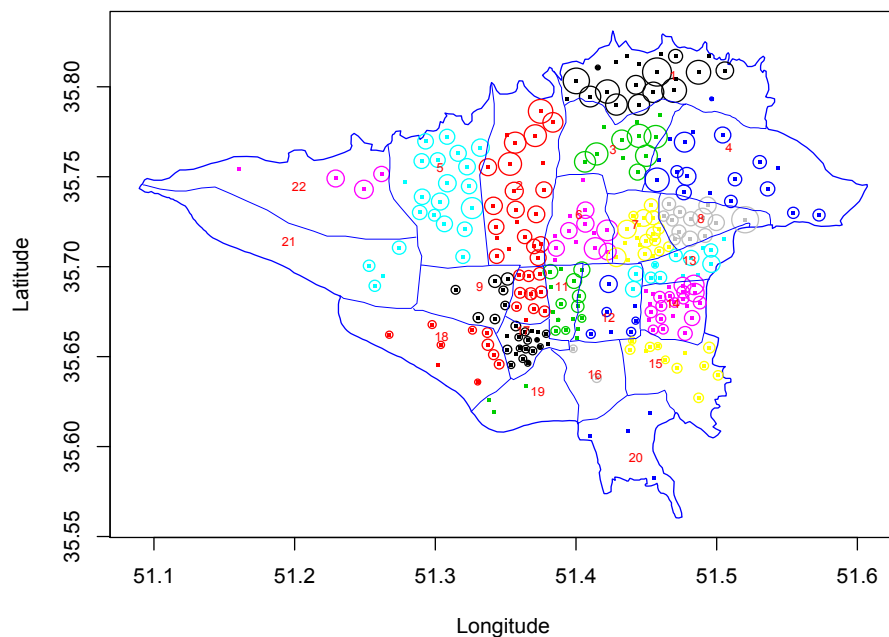
۵ داده‌های قیمت مسکن شهر تهران

در این قسمت داده‌های مربوط به متوسط قیمت مسکن در مناطق مختلف تهران بررسی می شود. چون داده‌های قیمت مسکن دارای همبستگی از نوع فضایی (مکانی) است، با استفاده از سه مدل رگرسیون فضایی مدل بندی می شود. داده‌های قیمت مسکن برای شهر تهران به صورت روزانه در سامانه اطلاعات بازار املاک ایران ثبت می شود که این داده‌ها را می توان در سایت www.hmi.mrud.ir مشاهده کرد. قیمت مسکن نوساز (تا یک سال ساخت) برای محله‌های ۲۲ منطقه در شهریور ماه سال ۱۳۹۴ استخراج شد. موقعیت مناطق و داده‌های در دسترس روی نقشه تهران در شکل ۱ مشخص شده است. دایره بزرگتر نشان دهنده قیمت مسکن بیشتر و اعداد منطقه را نشان می دهند. خلاصه آماره‌های مربوط به این مجموعه داده در جدول ۱ ارائه شده است. میانگین قیمت مسکن حدود چهار میلیون، ماکسیمم قیمت مسکن دوازده میلیون و مینیمم قیمت مسکن حدود یک و نیم میلیون است. نمودار بافت‌نگار، جعبه‌ای و نرمال چندک-چندک داده‌ها در شکل ۲ رسم شده است

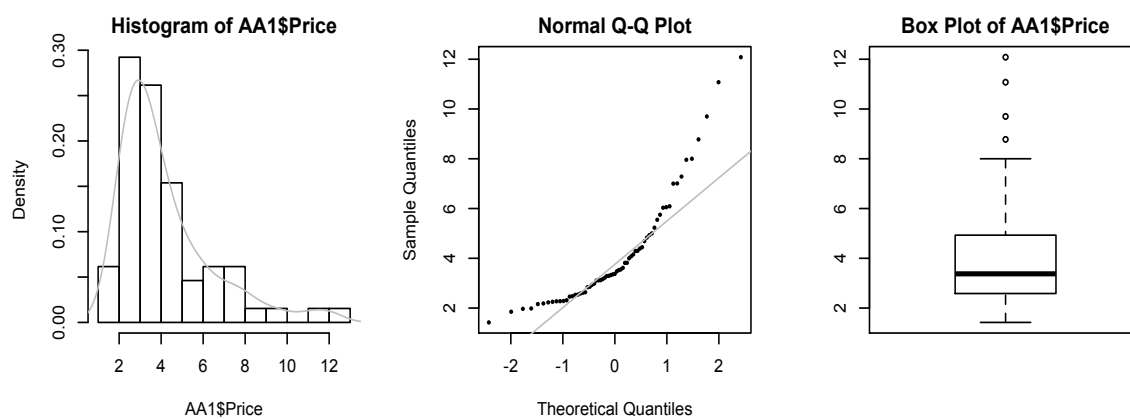
جدول ۱: آماره‌های توصیفی داده‌های مسکن تهران.

متغیر	میانگین	چارک اول	میانه	چارک سوم	مینیمم	ماکسیمم
قیمت مسکن	۱۷۷/۴	۵۸۴/۲	۳۷۶/۳	۹۳۰/۴	۴۲۳/۱	۰۸۰/۱۲

به دلیل عدم نرمال بودن داده‌ها از تبدیل لگاریتمی بر روی داده‌ها استفاده شد و مقدار آماره شاپیرو ویلک $0.963/0$ به دست آمد که فرض نرمال بودن پذیرفته می شود. مساحت ملک به عنوان متغیر کمکی در نظر گرفته شد. داده‌ها با سه مدل رگرسیون فضایی SEM، SLM و SDM داده‌ها با سه رهیافت ماکسیمم درستنمایی، بیزی و بیزی تقریبی تحلیل و نتایج بیزی تقریبی در جدول ۲ آورده شده است. برای تحلیل بیزی معمولی از الگوریتم گیبز و متروپلیس هاستینگز با ۵۰۰۰۰ تکرار



شکل ۱: موقعیت داده‌های متوسط قیمت مسکن شهر تهران.



شکل ۲: نمودار جعبه‌ای، چنک-چنک و هیستوگرام داده‌های مسکن تهران.

استفاده شد. برازش و مقادیر مجذور میانگین توان دوم خطای نسبی^۹ (RRMSE) به صورت

$$RRMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n \frac{(y_i - \hat{y}_i)^2}{\hat{y}_i^2}},$$

برای سه مدل SLM، SDM و SEM محاسبه و در جدول ۳ ارائه شده است. نتایج جدول ۳ بیانگر برتری مدل SLM نسبت به دو مدل SEM و SDM و دقت بهتر رهیافت بیزی و بیزی تقریبی نسبت به درست‌نمایی و معادل بودن نتایج رهیافت بیزی معمولی با بیزی تقریبی برای این مجموعه داده است، به طوری که زمان اجرای رهیافت بیزی تقریبی کمتر از یک دقیقه

^۹Relative Root Mean Square Error

جدول ۲: برآورد بیزی تقریبی پارامترها و خلاصه آماره‌ها.

مدل	پارامتر	برآورد	انحراف معیار	چندک ۰/۰۲۵	چندک ۰/۵	چندک ۰/۹۷۵
SLM	ρ	۰/۸۱۷۳	۰/۰۸۲۳	۰/۶۳۸۶	۰/۸۲۹۲	۰/۹۵۷۸
	β	-۰/۴۱۰۲	۰/۱۳۰۰	-۰/۶۶۴۱	-۰/۴۱۱۲	-۰/۱۳۱۷
	β_{Area}	۰/۰۰۶۹	۰/۰۰۱۱	۰/۰۰۴۶	۰/۰۰۶۹	۰/۰۰۹۳
SDM	ρ	۰/۸۲۶۱	۲/۵۸۴	۳/۳۷۶	۴/۹۳۰	۱/۴۲۳
	β	-۰/۳۴۴۹	۰/۲۴۶۶	-۰/۸۸۹۷	-۰/۳۴۲۳	۰/۱۹۹۸
	β_{Area}	۰/۰۰۶۸	۰/۰۰۱۱	۰/۰۰۴۵	۰/۰۰۶۹	۰/۰۰۹۴
	γ	-۰/۰۰۰۷	۰/۰۰۲۶	-۰/۰۰۷۴	-۰/۰۰۰۸	۰/۰۰۶۱
SEM	λ	۰/۹۲۰۲	۰/۰۵۸۶	۰/۷۷۳۱	۰/۹۴۱۱	۰/۹۸۹۷
	β	۰/۵۵۹۸	۰/۳۱۹۲	-۰/۹۹۰۱	۰/۵۹۶۴	۱/۸۹۷۳
	β_{Area}	۰/۰۰۶۴	۰/۰۰۱۰	۰/۰۰۴۸	۰/۰۰۶۵	۰/۰۰۸۳

جدول ۳: مقادیر مجذور میانگین مربع خطای نسبی برای سه مدل.

رهیافت	SLM	SEM	SDM
ماکسیمم درست‌نمایی	۰/۲۴۱۸	۰/۲۴۴۷	۰/۲۴۳۵
بیزی معمولی	۰/۲۳۶۷	۰/۲۴۰۲	۰/۲۳۸۹
بیزی تقریبی	۰/۲۳۷۱	۰/۲۳۹۷	۰/۲۳۸۸

و زمان اجرای رهیافت بیزی معمولی بیش از یک ساعت به طول انجامید.

بحث و نتیجه‌گیری

در این مقاله سه مدل رگرسیونی فضایی در اقتصادسنجی برای مشاهداتی که دارای وابستگی مکانی هستند به‌کار گرفته شد. نحوه برآورد پارامترها با استفاده از روش ماکسیمم درست‌نمایی و رهیافت بیزی تقریبی ارائه شد. در یک مجموعه داده واقعی به بررسی ارتباط قیمت و مساحت مسکن در شهریور ماه ۱۳۹۴ شهر تهران پرداخته شد. با توجه به عدم نرمال بودن داده‌ها از تبدیل لگاریتمی داده‌ها استفاده شد. پس از برازش مدل‌ها و مقایسات انجام شده مدل مناسب مدل SLM انتخاب و در هر سه مدل و با هر دو رهیافت بیزی تقریبی و درست‌نمایی وجود اثر فضایی و ارتباط مستقیم بین قیمت مسکن و مساحت پذیرفته می‌شود. با استفاده از معیار مجذور میانگین مربع خطای نسبی و براساس این دو مجموعه داده مشاهده شد که رهیافت بیزی تقریبی برای این نوع از مدل‌های اقتصادسنجی فضایی نسبت به رهیافت درست‌نمایی از عملکرد تقریباً بهتری برخوردار است و زمان محاسبات رهیافت بیزی تقریبی کمتر از رهیافت بیزی معمولی است. در مدل‌های معرفی شده برای راحتی محاسبات معمولاً توزیع نرمال برای خطاها منظور می‌شود، به عنوان پیشنهاد می‌توان در مدل‌های اقتصادسنجی فضایی از توزیع‌های انعطاف‌پذیرتر مثل چوله نرمال، چوله نرمال بسته و چوله تی استفاده کرد که کلاس بزرگتری نسبت به توزیع نرمال هستند.

مراجع

- Anselin, L. (1988), *Spatial Econometrics: Method and Models In: Studies in Operational Regional Science*, Springer.
- Anselin, L. (2010), Thirty Years of Spatial Econometrics, *Papers in Regional Science*, **89**, 3-25.
- Arbia, G. (2016), Spatial Econometrics: A Broad View, *Foundations and Trends in Econometrics*, **8**, 145-265.
- Barry, R. P., Pace, R. K. (1997), Kriging with Large Data Sets using Sparse Matrix Techniques, *Communication Statistics: Computation and Simulation*, **26**, 619-629.
- Bivand, R.s., Gómez-Rubiob, V., Rue, H. (2014), Approximate Bayesian Inference for Spatial Econometrics Models, *Spatial Statistics*, **9**, 146–165.
- Davidson, R., Mackinnon, J. G. (1993), *Estimation and Inference in Econometrics*, Oxford University, New York.
- Eidsvik, J., Finley, A. O., Banerjee, S., Rue, H. (2012), Approximate Bayesian Inference for Large Spatial Datasets Using Predictive Process Models, *Computational Statistics and Data Analysis*. **56**, 1362-1380.
- Kelejian, H., Piras, G. (2017), *Spatial Econometrics*. London, England: Academic Press.
- LeSage, J. (1999), *Spatial Econometrics*, Department of Economics University of Toledo, .
- LeSage, J., Pace, R. K. (2009), Introduction to Spatial Econometrics, *Chapman and Hall/CRC*.
- LeSage, J., Pace, R. K., Lam, N., Campanella, R., Liu, X. (2011), New Orlean Business Recovery in the Aftermath of Hurricane Katrina, *Journal of the Royal Statistical Society, Ser. A*, **174**, 1007-1027.
- Lesage, J., Ficher, M. (2008), Spatial Growth Regression: Model Specification, Estimation and Interpretation, *Spatial Economic Analysis*, **3**, 275–304.
- Rue, H., Held, L. (2005), *Gaussian Markov Random Fields. Theory and Applications*, Chapman and Hall, London.
- Rue, H. Martino, S., Chopin, N. (2009), Approximate Bayesian Inference for Latent Gaussian Model by Using Integrated Nested Laplace Approximations (with discussion), *Journal of the Royal Statistical Society*. **71**, 319–392.



مقدمه‌ای بر تقریب لاپلاس آشیانه‌ای جمع بسته: رهیافت سریع و دقیق بیزی

سحر خلفی^۱، علی محمدیان مصمم، علی آقامحمدی
گروه آمار، دانشگاه زنجان

چکیده: یکی از عملیات‌های مهم در استنباط بیزی محاسبه انتگرال‌هایی با بعد بالاتر است. روش‌های قدیمی مانند روش لاپلاس که با بسط مرتبه دوم تیلور از راه تجزیه انتگرال محاسبه می‌شود، وجود دارد که بسیار وقت‌گیر و اغلب مشکل همگرایی دارد. با بسط و توسعه یک روش تودرتو از یک ایده کلاسیک با روش‌های عددی جدید برای ماتریس‌های تنک در این مقاله شیوه‌ای از تقریب لاپلاس آشیانه‌ای جمع بسته ($INLA$) در انجام تقریب استنباط بیزی برای مدل‌های گاوسی پنهان مورد مطالعه قرار می‌گیرد. همچنین در این مقاله به تشریح خروجی‌های $INLA$ و سرعت و دقت بسته $R-INLA$ می‌پردازیم.

واژه‌های کلیدی: تقریب لاپلاس آشیانه‌ای جمع بسته، تقریب لاپلاس، ماتریس‌های تنک، مدل‌های گاوسی پنهان، استنباط

بیزی

کد موضوع بندی ریاضی (۲۰۱۰): 62M30.

۱ مقدمه

شبیه‌سازی مبنی بر استنباط از طریق الگوریتم زنجیره مارکوف مونت کارلو ($MCMC$) در اوایل دهه ۱۹۹۰ به جامعه آماری معرفی شد که معمولاً برای محاسبه‌ی توزیع‌های پسینی مفید است. اما وجود پارامترهای زیاد در ساختار سلسله مراتبی این مدل‌ها مخصوصاً زمانی که بعد میدان پنهان بزرگ باشد و وجود همبستگی بین عناصر میدان پنهان ویژگی‌های همگرایی و آمیختگی این الگوریتم را با مشکل مواجه می‌کند. از طرفی به دلیل پیچیده بودن با مشکلی مانند طولانی بودن زمان محاسبات مواجه می‌شود که برای حل این مشکلات **رو و همکاران (۲۰۱۷)** روش تقریب لاپلاس آشیانه‌ای جمع بسته ($INLA$) که ترکیبی از تقریب لاپلاس و انتگرال‌های عددی است را معرفی کردند و نشان دادند که این روش در مدت زمان بسیار کوتاه و با حفظ دقت برآمد پارامترها تقریبی قابل قبول ارائه می‌دهد. در واقع $INLA$ برای برازش مدل‌های متفاوتی به ویژه داده‌های فضایی در زمینه‌های بیزی کاربرد دارد و خروجی آن شامل توزیع‌های پسینی است که می‌توان از راه میانگین، واریانس و

^۱ نام و ایمیل ارائه دهنده مقاله: سحر خلفی، s.khalafi@znu.ac.ir

چندک مورد ارزیابی قرار گیرد. نمونه‌هایی که از $INLA$ برای تجزیه و تحلیل آماری اخیراً استفاده شده عبارت هستند از: تطبیق پذیری پیشینی وزنی در رگرسیون تعمیم یافته (هلد و ساتر، ۲۰۱۶)، سرایت مالاریا در آفریقا (نور و همکاران، ۲۰۱۴)، مطالعات ارتباط بین دسترسی به مسکن و سلامت و رفاه در کلان شهرها (کانت و همکاران، ۲۰۱۶)، متآنالیز شبکه (ساتر و همکاران، ۲۰۱۵) و غیره.

سه مولفه اصلی $INLA$ در تقریب استنباط بیزی شامل مدل‌های گاوسی پنهان (LGM)، میدان تصادفی گاوسی مارکف ($GMRf$) و تقریب لاپلاس هستند که در بخش ۲ معرفی خواهند شد. در بخش ۳ تقریب لاپلاس آشیانه‌ای جمع بسته تشریح می‌شود و در بخش ۴ با یک مثال کاربردی ساده به توصیف بیشتر این روش پرداخته می‌شود. بخش ۵ شامل بحث و نتیجه گیری است.

۲ مولفه‌های اصلی $INLA$

علیرغم سادگی ساختار مدل‌های گاوسی پنهان اغلب مدل‌های مورد استفاده از آمار و سایر رشته‌های کاربردی نظیر مدل‌های خطی تعمیم یافته، مدل‌های جمعی تعمیم یافته، اسپلاین‌ها، مدل‌های رگرسیون شبه پارامتری، مدل‌های فضایی و فضایی-زمانی را می‌توان در قالب مدل‌های گاوسی پنهان در نظر گرفت. مدل‌های گاوسی پنهان را با استفاده از مدل سلسله مراتبی سه مرحله‌ای زیر که در آن مشاهدات y را به طور مشروط مستقل و θ_1 را ابر پارامتر و x را متغیر پنهان می‌توان فرض کرد را به دست آورد.

$$y|x, \theta_1 \sim \prod_{i \in \tau} \pi(y_i|x_i, \theta_1)$$

که در آن میدان تصادفی گاوسی پنهان x به صورت:

$$x|\theta_2 \sim N(\mu(\theta_2), Q^{-1}(\theta_2))$$

می‌باشد و θ_2 پارامتر و Q ماتریس دقت و $\theta = (\theta_1, \theta_2)$ ابر پارامتر می‌باشد. توزیع پسین به صورت:

$$\pi(x, \theta|y) \propto \pi(\theta)\pi(x|\theta) \prod_{i \in \tau} \pi(y_i|x_i, \theta)$$

می‌باشد. از آنجایی که این توزیع اغلب فرم بسته‌ای ندارد نیازمند روش‌های عددی می‌باشیم. علاوه بر این فرضیه‌های مهم زیر نیز در نظر گرفته می‌شوند.

۱. تعداد ابرپارامترها $|\theta|$ کم و به طور نمونه بین ۲ تا ۵ می‌باشد.

۲. توزیع میدان پنهان $x|\theta$ گاوسی می‌باشد و وقتی که بعد n بزرگ باشد (10^3 تا 10^5) به یک $GMRf$ نیاز داریم.

۳. داده‌های y به طور مشروط مستقل از θ, x هستند و اشاره دارد که هر مشاهده y_i تنها به یک مولفه از میدان پنهان x_i بستگی دارد.

x یک $GMRf$ است اگر یک میدان تصادفی گاوسی با خاصیت استقلال شرطی باشد. به این معنی که x_i, x_j به شرط سایر داده‌های باقیمانده x_{-ij} مستقل باشند. یک مثال ساده مدل اتو رگرسیو مرتبه اول به صورت زیر است:

$$x_t = \phi x_{t-1} + \varepsilon_t; \quad t = 1, \dots, m$$

که در آن ε فرآیند اغتشاش و $|\phi| < 1$ و اندیس t نشان دهنده‌ی زمان است. به سادگی می‌توان نشان داد که در این مدل همبستگی بین x_t, x_s ، $|s-t| > 1$ ، داده‌های مستقل شرطی روی x_{-st} هستند. بنابراین ماتریس دقت مربوطه یک ماتریس سه قطری و در نتیجه ماتریس تنک خواهد شد که باعث سادگی محاسبات می‌شود. در حالت کلی هدف تقریب انتگرال زیر وقتی $n \rightarrow \infty$ است:

$$I_n = \int_x \exp(nf(x)) dx$$

ایده اصلی این است که با استفاده از بسط تیلور $f(x)$ و تفسیر $nf(x)$ به عنوان مجموع لگ-درستنمایی و در نظر گرفتن x به عنوان پارامتر نامعلوم طبق قضیه حد مرکزی با درستنمایی توزیع گاوسی تقریب زده شود. این تقریب بویژه برای انتگرال‌های با ابعاد بالا بسیار مفید است. به عنوان مثال فرض کنید که علاقه‌مند به محاسبه توزیع حاشیه‌ای $\pi(\gamma_1)$ از توزیع توأم $\pi(\gamma)$ به صورت

$$\pi(\gamma_1) = \frac{\pi(\gamma)}{\pi(\gamma_{-1}|\gamma_1)} \approx \frac{\pi(\gamma)}{\pi_G(\gamma_{-1}; \mu(\gamma_1), Q(\gamma_1))} \Big|_{\gamma_{-1}=\mu(\gamma_1)}$$

هستیم، که در آن $\pi(\gamma_{-1}|\gamma_1)$ با استفاده از توزیع گاوسی تقریب زده شده است که در مدل گاوسی پنهان بخش قبل در واقع $\gamma = (x, \bullet)$ است.

۳ تقریب لاپلاس آشیانه‌ای جمع بسته

حال به معرفی $INLA$ برای حالتی که مدل گاوسی پنهان شامل پارامتر با بعد پایین و $x|\theta$ یک $GMRF$ است می‌پردازیم. همان‌گونه که اشاره شد هدف استنباط بیزی تقریب حاشیه‌های پسین به صورت

$$\pi(\theta_j|y); j = 1, \dots, |\theta|, \quad (x_i|y); i = 1, \dots, n$$

است. در این روش درستنمایی مستقل شرطی است به معنی اینکه y_i تنها به x_i, θ بستگی دارد. مدل کلی زیر را در نظر بگیرید:

$$\eta_i = g(\beta)u_{j(i)}$$

که در آن برای نمونه $i = 1, \dots, n$ ، $y_i|\eta_i \sim \text{poisson}(\exp(\eta_i))$ ، $y_i|\eta_i \sim N(\bullet, 1)$ ، $\beta \sim N(\bullet, 1)$ ، تابع یکنوا خوش رفتار و $u \sim N(\bullet, Q^{-1})$ است. اندیس $j(i)$ به معنی نگاشتی است که بعد u ثابت بوده و به n بستگی ندارد. محاسبه حاشیه‌های پسین برای β, u_j از آنجایی که یک حاصل ضرب گاوسی و یک غیر گاوسی داریم مشکل است. لذا تقریب به زیر مجموعه‌های کوچکتر تقسیم و از تقریب لاپلاس استفاده می‌شود:

$$\pi(\beta|y) \propto \pi(\beta) \int \prod_{i=1}^n \pi(y_i|\lambda_i = \exp(g(\beta)u_{j(i)})) \times \pi(u) du$$

برای هر u_j توزیع حاشیه‌ای به صورت

$$\pi(u_j|y) = \int \pi(u_j|\beta, y) \times \pi(\beta|y) d\beta$$

است و از آنجایی که β یک پارامتر تک بعدی است این محاسبات به سادگی انجام می‌گیرد.

$$\pi(u|\beta, y) \propto \prod_{i=1}^n \pi(y_i|\lambda_i = \exp(g(\beta)u_{j(i)})) \times \pi(u)$$

به همین ترتیب که در اینجاست چون u_j ثابت در نظر گرفته می شود محاسبه $\pi(u_j|\beta, y)$ شامل انتگرال گیری عددی است. حال هدف محاسبه توزیع پسینی برای هر θ_j است. برای این منظور یک تقریب برای

$$\pi(\theta|y) \propto \frac{\pi(\theta)\pi(x|\theta)\pi(y|x, \theta)}{\pi(x|\theta, y)}$$

در نظر گرفته می شود که در آن مخرج کسر با یک تقریب گاوسی به صورت

$$\begin{aligned}\pi(x|y, \theta) &\propto \exp\left(-\frac{1}{2}x^T Q(\theta)x + \sum_i \log \pi(y_i|x_i, \theta)\right) \\ &= (2\pi)^{-n/2} |P(\theta)|^{1/2} \exp\left(-\frac{1}{2}(x - \mu(\theta))^T P(\theta)(x - \mu(\theta))\right)\end{aligned}$$

محاسبه می شود، که در آن $P(\theta) = Q(\theta) + \text{diag}(c(\theta))$ ، $\mu(\theta)$ مکان مد، $c(\theta)$ حاوی منفی مشتقات دوم نسبت به x_i از لگ-درستمایی در نقطه مد است. مزیت این تقریب این است که مسئله به یک GMR خلاصه شده، در نتیجه محاسبات ساده تر و تقریب دقیق تر خواهد شد. سرانجام توزیع های پسینی برای میدان پنهان را می توان تقریب زد. همانند قبل در این حالت داریم:

$$\pi(x_i|y) = \int \pi(x_i|\theta, y)\pi(\theta|y)d\theta$$

در اینجا ما به انتگرال گیری روی $\pi(\theta|y)$ نیاز داریم اما هزینه محاسبات از انتگرال گیری عددی استاندارد از مرتبه نهایی بسته به بعد θ است. همچنین ما به تقریب $\pi(x_i|\theta, y)$ برای تمام $i = 1, \dots, n$ با n بسیار بزرگ نیاز داریم که با استفاده از تقریب لاپلاس امکان پذیر است.

۴ مثال کاربردی

در این بخش با یک مثال شبیه سازی به معرفی بسته $R - INLA$ در نرم افزار R می پردازیم. برای این منظور فرض کنید که داده ها از

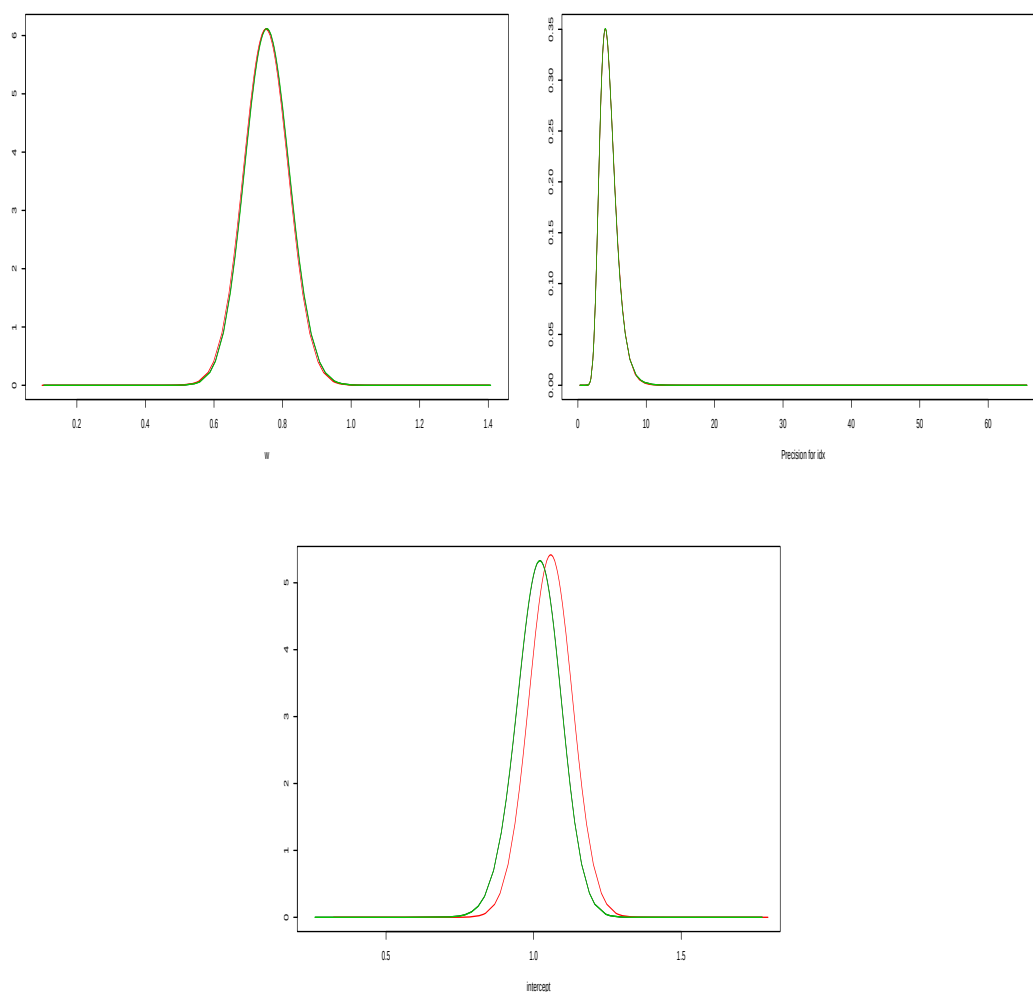
$$y|\eta \sim \text{Poisson}(\exp(\eta))$$

که در آن $\eta_i = \mu + \beta w_i + u_{j(i)}$ ؛ $i = 1, \dots, n$ است، w متغیر کمکی، $u \sim N_m(0, \tau^{-1}I)$ می باشند، تولید شده باشند. برای تولید $n = 200$ داده ما فرض کرده ایم که $\mu = 1$ ، $\beta = 3/4$ و $\tau = 4/16$ باشد. حال با استفاده از پکیج $R - INLA$ به تحلیل بیزی داده های شبیه سازی شده می پردازیم.

شکل ۱ برآورد پسینی حاشیه ای اثرات ثابت $(intercept, w)$ و ابر پارامتر $(precision \text{ for } idx)$ را با استفاده از سه روش تقریب لاپلاس (خط مشکی)، تقریب گاوسی (خط قرمز) که مانع از انتگرال گیری روی θ شده و انتگرال دقیق روی θ (خط سبز) نشان می دهد. همچنین خلاصه آماری مربوط به برآورد پارامترها در جداول ۱-۳ ارائه شده است. همان گونه که ملاحظه می شود تقریب لاپلاس بسیار نزدیک به تقریب دقیق انتگرال می باشد.

۵ بحث و نتیجه گیری

محاسبات پیچیده تحلیل بیزی یکی از مشکلات اساسی کاربران روش های بیزی به ویژه وقتی که اندازه نمونه و بعد پارامترها زیاد است می باشد. در این مقاله روش $INLA$ به عنوان روش جایگزین معرفی گردید که می تواند بسیار سریع تر و حتی دقیق تر از روش های کلاسیک نظیر $MCMC$ باشد.



شکل ۱: برآورد پسینی حاشیه‌ای اثرات ثابت و ابر پارامتر با سه روش تقریب لاپلاس (خط مشکی)، تقریب گاوسی (خط قرمز) و انتگرال دقیق روی θ (خط سبز)

جدول ۱: خلاصه آماری مربوط به برآورد پارامترها به روش تقریب لاپلاس

مد	چارک سوم	میانه	چارک اول	انحراف معیار	میانگین	
۱/۰۲۱	۱/۱۶۲	۱/۰۱۹	۰/۸۶۴	۰/۰۷۶	۱/۰۱۷	عرض از مبدأ
۰/۷۵۳	۰/۸۸۳	۰/۷۵۴	۰/۶۲۷	۰/۰۶۵	۰/۷۵۴	w
۴/۰۱۰	۷/۶۰۰	۴/۳۵۰	۲/۵۷۰	۱/۲۹۰	۴/۵۴۰	دقت

مراجع

Held, L., and Sauter, R. (2017), Adaptive Prior Weighting in Generalized Regression, *Biometrics*, **73**, 242-251.

Kandt, J., Chang, S. S., Yip, P., and Burdett, R. (2017), The Spatial Pattern of Premature Mortality

جدول ۲: خلاصه آماری مربوط به برآورد پارامترها به روش تقریب گاوسی

مد	چارک سوم	میانه	چارک اول	انحراف معیار	میانگین	
۱/۰۵۸	۱/۲۰۲	۱/۰۵۸	۰/۹۱۳	۰/۰۷۴	۱/۰۵۸	عرض از مبدأ
۰/۷۵۱	۰/۸۷۹	۰/۷۵۱	۰/۶۲۳	۰/۰۶۵	۰/۷۵۱	w
۴/۰۲۰	۷/۵۳۰	۴/۳۶۰	۲/۵۷۰	۱/۲۷۰	۴/۵۴۰	دقت

جدول ۳: خلاصه آماری مربوط به برآورد پارامترها به روش انتگرال دقیق

مد	چارک سوم	میانه	چارک اول	انحراف معیار	میانگین	
۱/۰۲۱	۱/۱۶۲	۱/۰۱۸	۰/۸۶۵	۰/۰۷۶	۱/۰۱۷	عرض از مبدأ
۰/۷۵۴	۰/۸۸۴	۰/۷۵۵	۰/۶۲۸	۰/۰۶۵	۰/۷۵۵	w
۴/۰۳۰	۷/۶۴۰	۴/۳۵۰	۲/۵۶۰	۱/۳۰۰	۴/۵۳۰	دقت

in Hong Kong: How Does it Relate to Public Housing? *Urban Studies*, **54**, 1211-1234.

Noor, A. M., Kinyoki, D. K., Mundia, C. W., Kabaria, C. W., Mutua, J. W., Alegana, V. A., Fall, I. S., and Snow, R. W. (2014), The Changing Risk of Plasmodium Falciparum Malaria Infection in Africa: 2000-10: A Spatial and Temporal Analysis of Transmission Intensity, *The Lancet*, **383**, 1739-1747.

Rue, H., Riebler, A., Sørbye, S.H., Illian, J.B., Simpson, D.P. and Lindgren, F.K., (2017), Bayesian Computing with INLA: A Review, *Annual Review of Statistics and Its Application*, **4**, 395-421.

Sauter, R. and Held, L. (2015), Network Meta-Analysis with Integrated Nested Laplace Approximations, *Biometrical Journal*, **57**, 1038-1050.



میانگین‌گیری بیزی مدل‌های رگرسیونی فضایی تحت توزیع گاوسی وارون

زهرا رحیمیان آزاد^۱، افشین فلاح، زهرا زارع‌پور یکنمی
گروه آمار، دانشگاه بین‌المللی امام خمینی

چکیده: میانگین‌گیری بیزی مدل‌ها روشی توانمند در مدل‌سازی است که برخلاف روش مرسوم گزینش مدل، از اطلاعات همه یا مجموعه‌ای از بهترین مدل‌ها استفاده می‌کند. از جمله برتری‌های این روش بر سایر روش‌های مدل‌سازی می‌توان به توانایی آن در لحاظ نمودن عدم حتمیت واقعی فرایند مدل‌سازی اشاره کرد. این مقاله میانگین‌گیری بیزی مدل‌های رگرسیونی فضایی در شرایطی که متغیر پاسخ دارای توزیع گاوسی وارون است را مورد توجه قرار می‌دهد. رویکرد پیشنهادی در قالب یک مطالعه شبیه‌سازی مورد ارزیابی قرار گرفته است.

واژه‌های کلیدی: داده‌های فضایی، میانگین‌گیری بیزی مدل‌ها، توزیع گاوسی وارون.
کد موضوع‌بندی ریاضی (۲۰۱۰): 62H11, 62F15, 62J12

۱ مقدمه

روش‌های کلاسیک مدل‌سازی آماری برپایه فرض استقلال مشاهدات توسعه یافته‌اند. این در حالی است که در بسیاری از کاربردها مشاهدات به دلیل موقعیت قرار گرفتن در فضای مورد مطالعه مستقل نیستند. تحلیل چنین داده‌هایی که داده‌های فضایی نامیده می‌شوند، با استفاده از روش‌های کلاسیک به دلیل در نظر نگرفتن وابستگی مشاهدات می‌تواند به صورت جدی گمراه‌کننده باشد. در داده‌های فضایی وابستگی مشاهدات، تابعی از فاصله مشاهدات از یکدیگر است. به این معنی که معمولاً مشاهدات نزدیک‌تر وابستگی بیشتری دارند. مجموعه روش‌های آماری که وابستگی مشاهدات را در مدل‌سازی و استنباط مد نظر قرار می‌دهند، روش‌های آمار فضایی نامیده می‌شوند. امروزه مباحث آمار فضایی در علوم مختلف از جمله زمین‌شناسی (ماترون، ۱۹۶۲)، همه‌گیرشناسی (دیگل، ۱۹۹۳)، کشاورزی (اسمیت، ۱۹۳۸)، شیمی (ویتن، ۱۹۶۲) و غیره کاربرد گسترده‌ای دارند. کرسی (۱۹۹۳) یک منبع کامل درباره روش‌های مختلف آمار فضایی و کاربردهای آن است. از طرفی در شیوه مرسوم مدل‌سازی که گزینش مدل نامیده می‌شود، تحلیل‌گر براساس یک یا چند معیار، مدلی را از میان

^۱ نام و ایمیل ارائه دهنده مقاله: زهرا رحیمیان آزاد، z.rahimian@edu.ikiu.ac.ir

مجموعه مدل‌های ممکن برمی‌گزینند. این شیوه مدل‌سازی از جنبه‌های مختلف مورد استفاده قرار می‌گیرد. از جمله اینکه معیار ارزیابی تابع نظر تحلیل‌گر است و از این رو ممکن است تحلیل‌گران متفاوت، مدل‌های متفاوتی را به عنوان بهترین مدل انتخاب کنند. از طرف دیگر انتخاب یک مدل باعث از دست رفتن اطلاعات سایر مدل‌ها می‌شود. میانگین‌گیری بیزی مدل‌ها، روشی نوین در مدل‌سازی است که بدون تکیه بر یک مدل منتخب از اطلاعات همه مدل‌ها استفاده می‌نماید. این روش به هر مدل احتمالی اختصاص و از میانگین همه مدل‌ها یا تعدادی از بهترین آنها استفاده می‌کند. اولین کار در این زمینه توسط لیمر (۱۹۷۸) انجام گرفت. **مادیگان و رفتی (۱۹۹۴)** دو روش پایه‌ای برای اجرای این شیوه مدل‌سازی ارائه کردند. **رفتی (۱۹۹۵)** و **رفتی و همکاران (۱۹۹۳، ۱۹۹۷)** میانگین‌گیری بیزی را در چارچوب مدل‌های رگرسیونی مورد بررسی قرار دادند. **هوئینگ (۱۹۹۴)** و **هوئینگ و همکاران (۱۹۹۶، ۱۹۹۹)** به نحوه انجام محاسبات و چگونگی تعیین مولفه‌های لازم در اجرای این روش پرداختند. در این مقاله بر مبنای توزیع گاوسی وارون، نظریه لازم برای میانگین‌گیری بیزی مدل‌های رگرسیونی در شرایطی که مشاهدات دارای همبستگی فضایی هستند، توسعه داده شده است. در این راستا نحوه تعیین توزیع پیشین برای مدل‌ها و پارامترها، و چگونگی محاسبه توزیع پسین متناظر آنها، مورد بحث قرار گرفته است. در بخش ۲ روش میانگین‌گیری بیزی مدل‌ها به صورت اجمالی شرح داده می‌شود و در بخش سوم مدل پیشنهادی ارائه و چگونگی برازش آن مورد بحث قرار می‌گیرد. نحوه محاسبه کمیت‌های مورد نیاز برای روش میانگین‌گیری بیزی در این بخش توضیح داده می‌شود. در بخش ۴ روش پیشنهادی در قالب یک مطالعه شبیه‌سازی مورد ارزیابی قرار می‌گیرد.

۲ میانگین‌گیری بیزی مدل‌ها

فرض کنید $\mathcal{M} = \{M_1, \dots, M_T\}$ نشان‌دهنده فضای مدل، D مجموعه مشاهدات و Δ کمیت مورد علاقه مانند یک مشاهده در آینده یا یک پارامتر باشد. در این صورت، توزیع پسین Δ را می‌توان با استفاده از قاعده احتمال کل به صورت یک میانگین وزنی از توزیع‌های پیش‌بین تحت مدل‌های $M_1, \dots, M_T$ ، به شکل

$$\pi(\Delta|D) = \sum_{k=1}^T p(\Delta|M_k, D)\pi(M_k|D), \quad (1.2)$$

نوشت، که در آن وزن‌ها احتمال‌های پسین مدل‌ها هستند. برپایه رابطه (۱.۲) میانگین و واریانس پسین Δ به ترتیب به صورت

$$\begin{aligned} \mathbb{E}(\Delta|D) &= \sum_{k=1}^T \mathbb{E}(\Delta|M_k, D)\pi(M_k|D) \\ &= \sum_{k=1}^T \hat{\Delta}_k \pi(M_k|D), \end{aligned}$$

و

$$\begin{aligned} \text{Var}(\Delta|D) &= \mathbb{E}_M(\text{Var}(\Delta|D, M)) + \text{Var}_M(\mathbb{E}(\Delta|D, M)) \\ &= \sum_{k=1}^T (\text{Var}(\Delta|D, M_k) + \hat{\Delta}_k^2 \pi(M_k|D) - \mathbb{E}(\Delta|D)^2), \end{aligned} \quad (2.2)$$

هستند، که در آنها $\hat{\Delta}_k$ برآورد بیزی کمیت Δ را تحت مدل k ، $k = 1, \dots, T$ ، نشان می‌دهد (لیمر، ۱۹۷۸). در رابطه (۲.۲) مولفه $\text{Var}_M(\mathbb{E}(\Delta|D, M))$ عدم حتمیت بین مدل‌ها را نشان می‌دهد. در رابطه (۱.۲)، توزیع پیش‌بین کمیت مورد علاقه به صورت ۳

$$p(\Delta|M_k, D) = \int p(\Delta|\theta_k, M_k)\pi(\theta_k|M_k, D)d\theta_k, \quad (3.2)$$

قابل محاسبه است، که در آن θ_k بردار پارامترهای مدل M_k را نشان می‌دهد. احتمال پسین درست بودن مدل M_k با استفاده از قاعده بیز به صورت

$$\pi(M_k|D) = \frac{p(D|M_k)\pi(M_k)}{\sum_{j=1}^T p(D|M_j)\pi(M_j)}, \quad (4.2)$$

به دست می‌آید، که در آن

$$p(D|M_k) = \int p(D|\theta_k, M_k)\pi(\theta_k|M_k)d\theta_k, \quad (5.2)$$

درست‌نمایی حاشیه‌ای مدل M_k ، $\theta_k = (\beta_k, \lambda_k)$ بردار پارامترها در مدل M_k ، $p(D|\theta_k, M_k)$ تابع درست‌نمایی، $\pi(M_k)$ احتمال پیشین درست بودن مدل M_k و $\pi(\theta_k|M_k)$ توزیع پیشین پارامترهای مدل M_k است. برای محاسبه $\pi(\Delta|D)$ مولفه‌های تشکیل دهنده آن را مشخص و جایگزین کرد. بنابراین لازم است توزیع پیشین پارامترها، احتمال پیشین و پسین هر مدل و توزیع پیشین کمیت مورد علاقه را برای همه مدل‌های موجود در فضای مدل، مشخص کرد.

۳ مدل پیشنهادی

با فرض آنکه متغیر پاسخ دارای توزیع گاوسی واریان است، مدل رگرسیونی

$$y_i|\mathbf{x}_i^{(k)} \sim IG(\mu(\mathbf{x}_i^{(k)}), \lambda) \quad \mu^{-1}(\mathbf{x}_i^{(k)}) = \mathbf{x}_i'^{(k)}\beta^{(k)} + \phi_i^{(k)}, i = 1, \dots, n, \quad (1.3)$$

را در نظر بگیرید، که در آن $IG(\cdot, \cdot)$ توزیع گاوسی واریان را نشان می‌دهد، $\beta^{(k)} = (\beta_1^{(k)}, \beta_2^{(k)}, \dots, \beta_p^{(k)})'$ بردار ضرایب رگرسیونی، $\mathbf{x}_i^{(k)} = (x_{i1}^{(k)}, x_{i2}^{(k)}, \dots, x_{ip}^{(k)})'$ بردار متغیرهای تبیینی و $\phi^{(k)} = (\phi_1^{(k)}, \phi_2^{(k)}, \dots, \phi_n^{(k)})'$ اثرهای تصادفی فضایی است. اثر تصادفی فضایی $\phi_i^{(k)}$ را می‌توان به صورت شرطی بر اثرهای تصادفی فضایی همسایگی‌ها به صورت

$$\phi_i^{(k)}|\phi_{-i}^{(k)}, \sigma^2(k) \sim N\left(\xi \sum_{j \neq i} \frac{w_{ij}}{w_{i+}} \phi_j^{(k)}, \frac{\sigma^2(k)}{w_{i+}}\right),$$

در نظر گرفت، که در آن $\phi_{-i}^{(k)} = \{\phi_j^{(k)}; j \neq i\}$ ، $\sigma^2(k)$ واریانس مشترک اثرهای تصادفی فضایی، w_{ij} درایه (i, j) از ماتریس همسایگی است که به شکل

$$w_{ij} = \begin{cases} 1, & \text{اگر مشاهده‌ی } i \text{ و } j \text{ همسایه باشند.} \\ 0, & i = j \end{cases}$$

تعریف می‌شود. همچنین پارامتر ξ که در بازه $(\frac{1}{\lambda_{(1)}}, \frac{1}{\lambda_{(n)}})$ تغییر می‌کند، پارامتر هموارسازی فضایی است، که در آن $\lambda_{(1)} \geq \dots \geq \lambda_{(n)}$ مقادیر ویژه ماتریس $D_W^{-\frac{1}{2}} W D_W^{-\frac{1}{2}}$ هستند، $W = \left(\frac{w_{ij}}{w_{i+}}\right)$ ماتریس وزن و D_W ماتریس قطری با درایه‌های قطری $\{w_{1+}, \dots, w_{n+}\}$ است.

تابع درست‌نمایی مدل به صورت

$$L(\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k) | y_1, \dots, y_n) = \frac{\lambda^{\frac{n}{2}}}{(2\pi y_i^2)^{\frac{n}{2}}} \exp \left\{ -\frac{\lambda}{2} \sum_{i=1}^n \frac{(y_i - \mu_i^{(k)})^2}{y_i \mu_i^{(k)}} \right\}$$

$$\begin{aligned}
& \propto \lambda^{\frac{n}{2}} \exp \left\{ \frac{-\lambda}{2} \left((YX^{(k)}\beta^{(k)} - \mathbf{1}_n)' Y^{-1} (YX^{(k)}\beta^{(k)} - \mathbf{1}_n) \right. \right. \\
& \quad \left. \left. + \mathbf{2}\phi'^{(k)} YX^{(k)}\beta^{(k)} + \phi'^{(k)} Y\phi^{(k)} - \mathbf{2}\mathbf{1}_n'\phi^{(k)} \right) \right\} \\
& = \lambda^{\frac{n}{2}} \exp \left\{ \frac{-\lambda}{2} \left(Q_1(\beta^{(k)}) + g(\phi^{(k)}) \right) \right\} \quad (2.3)
\end{aligned}$$

نوشته می‌شود که در آن $Q_1(\beta^{(k)})$ یک نمایش درجه دوم برحسب $\beta^{(k)}$ ، Y یک ماتریس قطری با درایه‌های قطری $(y_1, \dots, y_n)$ و $g(\phi^{(k)})$ تابعی برحسب $\phi^{(k)}$ است.

۱.۳ تعیین توزیع پیشین پارامترها تحت مدل k ام

لازمه تحلیل بیزی مدل (۱.۳) آن است که برای پارامترهای آن، $(\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k))$ ، توزیع‌های پیشین مناسبی در نظر گرفته شود که بتوانند به خوبی باورهای پیشین تحلیل‌گر را در مدل‌سازی بازتاب دهند. از این رو، توزیع پیشین توام پارامترها را می‌توان با فرض استقلال پیشینی $(\beta^{(k)}, \lambda)$ و $(\phi^{(k)}, \sigma^2(k))$ به شکل

$$\begin{aligned}
\pi(\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k)) &= \pi(\beta^{(k)}, \lambda) \pi(\phi^{(k)}, \sigma^2(k)) \\
&= \pi(\beta^{(k)} | \lambda) \pi(\lambda) \pi(\phi^{(k)} | \sigma^2(k)) \pi(\sigma^2(k)), \quad (3.3)
\end{aligned}$$

نوشته شد. با توجه به این که دامنه تغییرات ضرایب رگرسیونی مجموعه اعداد حقیقی است و پارامتر λ ، تنها مقادیر مثبت را اختیار می‌کند، توزیع‌های پیشین تحت مدل M_k به صورت

$$\begin{aligned}
\beta^{(k)} | \lambda &\sim N_{p_k}(\eta_0, \lambda^{-1} V_0), \quad \beta^{(k)}, \eta_0 \in \mathbb{R}^{p_k} \\
\lambda &\sim \text{Gamma}\left(\frac{a_0}{2}, \frac{b_0}{2}\right), \quad \lambda, a_0, b_0 > 0,
\end{aligned}$$

در نظر گرفته می‌شوند، به گونه‌ای که $V_0 = \text{diag}(v_{01}, \dots, v_{0p_k})$ یک ماتریس قطری با درایه‌های قطری مثبت است و $\eta_0 = (\eta_{01}, \dots, \eta_{0p_k})$ و $\beta^{(k)} = (\beta_1^{(k)}, \dots, \beta_{p_k}^{(k)})$ ، کمیت‌های V_0 ، η_0 ، a_0 و b_0 که با اندیس صفر متمایز شده‌اند، ابرپارامترهای مدل بیزی را نشان می‌دهند و توسط تحلیل‌گر و براساس باورهای پیشین وی درباره پارامترها تعیین می‌شوند (مشکانی و همکاران، ۲۰۱۶). بنابراین توزیع پیشین توام $(\beta^{(k)}, \lambda)$ به صورت

$$\begin{aligned}
\pi(\beta^{(k)}, \lambda) &= \pi(\beta^{(k)} | \lambda) \pi(\lambda) \\
&\propto \lambda^{\frac{a_0 + p_k}{2} - 1} \exp \left\{ \frac{-\lambda}{2} b_0 \right\} \exp \left\{ \frac{-\lambda}{2} \left((\beta^{(k)} - \eta_0)' V_0^{-1} (\beta^{(k)} - \eta_0) \right) \right\}, \quad (4.3)
\end{aligned}$$

به دست می‌آید، که یک توزیع نرمال - گاما است. براساس لم بروک (بانرجی و همکاران، ۲۰۰۴) توزیع توام بردار اثرهای تصادفی فضایی، $\phi^{(k)}$ ، به شرط معلوم بودن $\sigma^2(k)$ به صورت

$$\phi^{(k)} | \sigma^2(k) \sim N_n \left(0, (D_W - \xi W)^{-1} \sigma^2(k) \right),$$

خواهد بود. از این رو اگر برای پارامتر $\sigma^2(k)$ نیز توزیع گامای وارون به شکل

$$\frac{1}{\sigma^2(k)} \sim \text{Gamma}(f_0, g_0),$$

در نظر گرفته شود، که در آنها f_0 و g_0 ابرپارامترهای توزیع پیشین هستند، توزیع توام $(\phi^{(k)}, \sigma^2(k))$ عبارت است از:

$$\begin{aligned}
\pi(\phi^{(k)}, \sigma^2(k)) &= \pi(\phi^{(k)} | \sigma^2(k)) \pi(\sigma^2(k)) \\
&\propto \exp \left\{ \frac{-1}{2\sigma^2(k)} \phi'^{(k)} (D_W - \xi W)^{-1} \phi^{(k)} \right\} \left(\frac{1}{\sigma^2(k)} \right)^{f_0 + 1} \exp \left\{ \frac{-g_0}{\sigma^2(k)} \right\} \quad (5.3)
\end{aligned}$$

۲.۳ محاسبه توزیع پسین توام پارامترها تحت مدل k ام

توزیع پسین توام پارامترها را می‌توان با استفاده از تابع درستنمایی (۲.۳) و توزیع پیشین توام (۴.۳) و (۵.۳)، به صورت

$$\begin{aligned} \pi(\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2)^{(k)} | y &\propto L(\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2)^{(k)} | y \pi(\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2)^{(k)} \\ &= \lambda^{\frac{a_0 + n + p_k}{2} - 1} \exp \left\{ -\frac{\lambda}{2} \left(b_0 + (\beta^{(k)} - \eta_0)' V^{-1} (\beta^{(k)} - \eta_0) \right) \right\} \\ &\times \exp \left\{ -\frac{\lambda}{2} (YX^{(k)} \beta^{(k)} - \mathbf{1}_n)' Y^{-1} (YX^{(k)} \beta^{(k)} - \mathbf{1}_n) \right\} \\ &\times \exp \left\{ \frac{-1}{2\sigma^2} \phi^{(k)'} (D_W - \xi W)^{-1} \phi^{(k)} \right\} \left(\frac{1}{\sigma^2} \right)^{f_0 + 1} \exp \left\{ \frac{-g_0}{\sigma^2} \right\} \\ &= \lambda^{\frac{a_0 + n + p_k}{2} - 1} \exp \left\{ -\frac{\lambda}{2} \left(Q_2(\beta^{(k)}) + u^{(k)} + g(\phi^{(k)}) \right) \right\} \\ &\times \left(\frac{1}{\sigma^2} \right)^{f_0 + 1} \exp \left\{ \frac{-1}{\sigma^2} \left(Q_2(\phi^{(k)}) \right) \right\} \end{aligned} \quad (6.3)$$

نوشت، که در آن

$$Q_2(\beta^{(k)}) = (\beta^{(k)} - \Sigma^{(k)-1} d^{(k)})' \Sigma^{(k)} (\beta^{(k)} - \Sigma^{(k)-1} d^{(k)}), \quad (7.3)$$

یک نمایش درجه دوم بر حسب $\beta^{(k)}$ را نمایش می‌دهد،

$$u^{(k)} = b_0 + \mathbf{1}_n' Y^{-1} \mathbf{1}_n + \eta_0' V^{-1} \eta_0 - d^{(k)'} \Sigma^{(k)-1} d^{(k)}, \quad (8.3)$$

کمیتی وابسته به مشاهدات است و

$$g(\phi^{(k)}) = 2\phi^{(k)'} YX^{(k)} \beta^{(k)} + \phi^{(k)'} Y \phi^{(k)} - 2\mathbf{1}_n' \phi^{(k)}, \quad (9.3)$$

$$Q_2(\phi^{(k)}) = \frac{\phi^{(k)'} (D_W - \xi W)^{-1} \phi^{(k)}}{2} - g_0, \quad (10.3)$$

$$\Sigma^{(k)} = (V^{-1} + X^{(k)'} YX^{(k)}),$$

$$d^{(k)} = V^{-1} \eta_0 + n\bar{x}^{(k)'}$$

همان‌گونه که ملاحظه می‌شود توزیع پسین توام پارامترها (۶.۳) پیچیده بوده و دارای نمایش بسته نیست. از این رو، برای استنباط بیزی لازم است از روش‌های مونت کارلو زنجیر مارکوفی برای نمونه‌گیری از توزیع پسین توام پارامترها و تقریب کمیت‌های پسینی مورد علاقه استفاده کرد. به منظور استفاده از الگوریتم گیبز برای نمونه‌گیری از توزیع پسین توام پارامترها، لازم است توزیع‌های پسین شرطی کامل پارامترها را به دست آورد.

۳.۳ توزیع‌های پسین شرطی کامل

برپایه‌ی توزیع پسین توام پارامترها در رابطه‌ی (۶.۳)، توزیع پسین شرطی کامل پارامترهای λ ، σ^2 ، $\beta^{(k)}$ و $\phi^{(k)}$ را می‌توان به ترتیب به شکل

$$\lambda | others \sim \text{Gamma} \left(\frac{a_0 + p_k + n}{2}, \frac{Q_2(\beta^{(k)}) + u^{(k)} + g(\phi^{(k)})}{2} \right),$$

$$\begin{aligned}\frac{1}{\sigma^2(k)} | others &\sim \text{Gamma} \left(f_0, Q_2(\phi^{(k)}) \right), \\ \pi(\beta^{(k)} | others) &\propto \prod_{i=1}^n \prod_{j=1}^r \exp \left\{ -\frac{\lambda}{\gamma} \left(Q_2(\beta^{(k)}) + g(\phi^{(k)}) \right) \right\}, \\ \pi(\phi^{(k)} | others) &\propto \prod_{i=1}^n \prod_{j=1}^r \exp \left\{ -\frac{\lambda}{\gamma} g(\phi^{(k)}) - \frac{1}{\sigma^2(k)} Q_2(\phi^{(k)}) \right\},\end{aligned}$$

محاسبه نمود، که $Q_2(\beta^{(k)})$ ، $u^{(k)}$ ، $Q_2(\phi^{(k)})$ و $g(\phi^{(k)})$ به ترتیب در روابط (۷.۳)، (۸.۳)، (۹.۳) و (۱۰.۳) معرفی شده‌اند. همان‌گونه که ملاحظه می‌شود این توزیع‌ها نمایش شناخته شده‌ای ندارند. به همین دلیل، برای نمونه‌گیری از آنها لازم است از الگوریتم متروپولیس-هاستینگز استفاده شود. از این‌رو، برای نمونه‌گیری از توزیع پسین توام پارامترها به تلفیقی از الگوریتم گیبز و متروپولیس-هاستینگز نیاز است.

۴.۳ محاسبه احتمال پسین مدل‌ها تحت مدل k ام

محاسبه احتمال پسین مدل‌ها، رابطه (۴.۲)، مستلزم محاسبه درستنمایی حاشیه‌ای و احتمال پیشین مدل‌ها است. برای این منظور می‌توان درستنمایی حاشیه‌ای، رابطه (۵.۲)، را برای مدل پیشنهادی با داشتن نمونه‌ای از توزیع پسین توام پارامترها که از روش مونت کارلو زنجیر مارکوفی به دست آمده به صورت

$$\begin{aligned}p(\mathbf{y}|M_k) &= \int p(\mathbf{y}|\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k), M_k) \pi(\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k) | M_k) d(\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k)) \\ &= \mathbb{E}_{\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k) | M_k} \left(p(\mathbf{y}|\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k), M_k) \right) \\ &\simeq \frac{1}{R} \sum_{r=1}^R p(\mathbf{y} | (\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k))^{(r)}, M_k)\end{aligned}\quad (۱۱.۳)$$

تقریب زد، که در آن $p(\mathbf{y}|\beta^{(k)}, \lambda, \phi^{(k)}, \sigma^2(k), M_k)$ تابع درستنمایی است. برای محاسبه احتمال پسین مدل M_k لازم است احتمال پیشین درست بودن برای همه مدل‌های موجود در فضای مدل $k = 1, \dots, T$ ، مشخص شوند. اگر هیچ اطلاع پیشینی در مورد مدل‌ها وجود نداشته باشد، توزیع مدل‌ها در فضای مدل، یکنواخت و به صورت

$$\pi(M_1) = \dots = \pi(M_T) = \frac{1}{T},$$

در نظر گرفته می‌شود. در این صورت احتمال پسین مدل M_k به صورت

$$\begin{aligned}\pi(M_k | \mathbf{y}) &= \frac{p(\mathbf{y}|M_k) \pi(M_k)}{\sum_{j=1}^T p(\mathbf{y}|M_j) \pi(M_j)} \\ &= \frac{p(\mathbf{y}|M_k)}{\sum_{j=1}^T p(\mathbf{y}|M_j)},\end{aligned}$$

نوشته می‌شود.

۵.۳ پنجره اوکام

با توجه به رابطه (۱۰.۲)، ملاحظه می‌شود که محاسبه مجموع مورد نظر به دلیل تعداد زیاد جملات آن دشوار است. به منظور کاهش تعداد مدل‌هایی که در مجموع (۱۰.۲) مشارکت دارند، از راهکاری تحت عنوان پنجره اوکام که توسط

مادیگان و رفتی (۱۹۹۴) پیشنهاد شده است، استفاده می‌شود. این راهکار از دو اصل پیروی می‌کند. بنابر اصل اول مدل‌هایی که در مقایسه با پیش‌بینی صورت گرفته براساس بهترین مدل‌ها، دارای پیش‌بینی نامطلوب‌تری هستند، کنار گذاشته می‌شوند. با اجرای این اصل مدل‌هایی که در مجموعه

$$\mathcal{A}' = \left\{ M_k : \frac{\max_{M_l \in \mathcal{M}} \pi(M_l | \mathbf{y})}{\pi(M_k | \mathbf{y})} \leq c \right\}$$

قرار نمی‌گیرند، از رابطه (۱.۲) خارج می‌شوند. آستانه c توسط تحلیل‌گر تعیین می‌شود. در اصل دوم که اصل تیغ اوکام نیز نامیده‌اند، مدل‌های پیچیده‌ای که نسبت به هم‌تایان ساده‌تر خود، کمتر از جانب داده‌ها حمایت می‌شوند نیز از مجموع کنار گذاشته می‌شوند. به عبارتی مدل‌هایی که متعلق به مجموعه

$$\mathcal{B} = \left\{ M_k : \exists M_l \in \mathcal{A} \text{ s.t. } M_l \subset M_k, \frac{\pi(M_l | \mathbf{y})}{\pi(M_k | \mathbf{y})} > 1 \right\}$$

هستند، از مجموع خارج می‌شوند. بنابراین رابطه (۱.۲)، به صورت

$$\pi(\Delta | \mathbf{y}) = \sum_{M_k \in \mathcal{A}} p(\Delta | M_k, \mathbf{y}) \pi(M_k | \mathbf{y})$$

بازنویسی می‌شود، که در آن $\mathcal{A} = \mathcal{A}' - \mathcal{B} \in \mathcal{M}$ ، مجموعه همه مدل‌هایی که در مجموع (۱.۲) قرار می‌گیرند.

۴ مطالعه شبیه‌سازی

در این بخش برای ارزیابی روش پیشنهادی از یک مطالعه شبیه‌سازی استفاده شده است. برای این منظور، یک متغیر پاسخ و هشت متغیر تبیینی در نظر گرفته شده است. در این صورت فضای مدل شامل $2^8 = 256$ مدل است. مدل اشباع شده که دربرگیرنده همه متغیرهای تبیینی است، به صورت

$$y_i | \mathbf{x}_i \sim IG(\mu(\mathbf{x}_i), \lambda), \quad \mu^{-1}(\mathbf{x}_i) = \beta_1 \mathbf{x}_1 + \beta_2 \mathbf{x}_2 + \dots + \beta_8 \mathbf{x}_8 + \phi_i, \quad i = 1, \dots, n, \quad (1.4)$$

نوشته می‌شود. مقادیر متغیرهای تبیینی به صورت مستقل و از توزیع نرمال شبیه‌سازی شده‌اند. سپس با فرض $\lambda = 2$ و

$$(\beta_1, \beta_2, \beta_3, \beta_4, \beta_5, \beta_6, \beta_7, \beta_8) = (1, 2, 3, 4, 5, 6, 7, 8),$$

مقادیر متغیرهای پاسخ از رابطه (۱.۴) شبیه‌سازی شده‌اند. روش پنجره اوکام چهارده مدل را براساس مقایسه احتمال پسین مدل‌ها به عنوان بهترین مدل‌های موجود معرفی می‌کند. برای مقایسه کارایی مدل فضایی پیشنهادی با مدل کلاسیک هم‌تای آن، از میانگین وزنی معیارهای اطلاع اکائیک، بیز و هانان - کوئین که به شکل

$$\begin{aligned} AIC &= -2l(\hat{\beta}) + 2k, \\ BIC &= -2l(\hat{\beta}) + k \log(n), \\ HQC &= -2l(\hat{\beta}) + 2k \log(\log n), \end{aligned} \quad (2.4)$$

تعریف می‌شوند و با احتمال پسین متناظرشان وزن‌دار شده‌اند، استفاده شده است. در رابطه (۲.۴)، k تعداد پارامترهای مدل، n اندازه نمونه، $\hat{\beta}$ برآورد بیز بردار ضرایب رگرسیونی تحت تابع زیان درجه دوم و $l(\cdot)$ تابع لگ درستیابی را نشان می‌دهند. مقادیر به دست آمده برای میانگین وزنی معیارهای یاد شده برای مدل فضایی پیشنهادی و هم‌تای کلاسیک آن در جدول ۱ نشان داده شده است. همان‌گونه که مشاهده می‌شود، مقادیر این میانگین‌ها برای روش میانگین‌گیری بیز مدل‌ها برای داده‌های فضایی از مقادیر متناظر برای میانگین‌گیری بیزی مدل‌ها کمتر هستند، که این موضوع نشان می‌دهد که مدل فضایی پیشنهادی به خوبی همبستگی‌های فضایی مشاهدات را در مدل‌بندی لحاظ می‌نماید.

جدول ۱: میانگین وزنی شاخص‌های نیکویی برازش برای مدل فضایی پیشنهادی و مدل کلاسیک.

HQC	BIC	AIC	میانگین‌گیری بیزی مدل‌ها
-۳۳۴/۹۳۲۷	-۳۳۱/۶۹۵۶	-۳۳۶/۷۶۶	مدل فضایی
-۳۲۷/۱۱۵	-۳۲۳/۹۸۲۲	-۳۲۸/۸۸۹۲	مدل کلاسیک

بحث و نتیجه‌گیری

روش‌های مرسوم مدل‌سازی آماری که برپایه مشاهدات مستقل بنا شده‌اند، در حالتی که مشاهدات دارای وابستگی فضایی هستند منجر به نتایج گمراه‌کننده می‌شوند. همچنین روش گزینش مدل با انتخاب یک مدل و مبنا قرار دادن آن در کاربردها، عدم حتمیت فرایند مدل‌سازی را به طور واقعی در نظر نمی‌گیرد. روش میانگین‌گیری بیزی مدل‌های رگرسیونی فضایی بدون از دست دادن اطلاعات مدل‌ها، از همه یا بهترین آنها بهره می‌گیرد و وابستگی فضایی داده‌ها را نیز لحاظ می‌کند.

مراجع

- Banerjee, S. Carlin, B. P. and Gelfand, A. E. (2004), *hierarchical Modeling and Analysis for Spatial Data*, Boca Raton: Chapman Hall/CRC.
- Cressie, N. (1993), *Statistics for Spatial Data*, John Wiley, New York.
- Diggle, P. J. (1993), Point Process Modeling in Environmental Epidemiology, In: Barnett, V. and Turkman, K. eds. *Statistics in Environment SPRUCE*, New York, Willey, Vol. 1.
- Fairfield, S. H. (1938), An Empirical Law Describing Heterogeneity in the Yields of Agricultural Crops, *Journal of Agricultural science*, **28**, 1-23.
- Hoeting J. (1994), Accounting for Uncertainty in Linear Regression Models, Ph.D Disertation, Department of Statistics, University of Washington.
- Hoeting, J., Raftery, A. and Madigan, D. (1996), A Method for Simultaneous Variable Selection and Outlier Identification in Linear Regression Models, *Journal of Statistical Computation and Simulation*, **22**, 251-271.
- Hoeting, J., Raftery, A. and Madigan, D. (1999), Bayesian Simultaneous Variable Selection and Transformation Selection in Linear Regression Models, Technical Report 9905, Dept. Statistics, Colorado State Univ. Available at www.colostate.edu.
- Leamer, E. E. (1978), *Specification Searches*, Wiley, New York.
- Madigan, D. and Raftery, A. (1994), Model Selection and Accounting for Model Uncertainty in Graphical Models Using Occam's Window, *Journal of the American Statistical Association*, **89**, 1535-1546.

- Matheron, G. (1962), *Traite de Geostatistique Appliquee*, Tome I. *Memoires du Bureau de Recherches et Minieres*, **14**, Editions: Technip, Paris.
- Meshkani, M. R. Fallah, A. and Kavousi, A. (2016), Bayesian analysis of covariance under inverse Gaussian model, *Journal of Applied Statistics*, **43**, 280-298.
- Raftery, A. E. (1995), Bayesian Model Selection in Social Research (With Discussion), in *Sociological Methodology 1995* (P. V. Marsden, ed.), 111-195. Blakwell. Cambridge, MA.
- Raftery, A. E., David, M. and Jenifer, A. H. (1993) Model Selection and Accounting for Model Uncertainty in Generalized Linear Models, Technical Report, **262**.
- Raftery, A. E., David, M. and Jenifer, A. H. (1997) Bayesian model averaging in Linear regression Models, *Journal of the American Statistical Association*, **97**, 179-191.
- Whitten, E. H. T. (1962), An reply to Chayes and Suzuki *Journal of Petrology*, **4**, 313-316.



مدل‌بندی توأم بیزی داده‌های فضایی-زمانی با گمشدگی غیرقابل چشم‌پوشی

سمیرا زحمتکش^۱، محسن محمدزاده
گروه آمار، دانشگاه تربیت مدرس

چکیده: اغلب داده‌های فضایی-زمانی در محیط‌های خارج از آزمایشگاه جمع‌آوری می‌شوند، به طوری که برخی عوامل تصادفی یا غیرتصادفی، بروز گمشدگی در داده‌ها را اجتناب ناپذیر می‌نمایند. از طرفی به دلیل وجود وابستگی بین مشاهدات، مقادیر گمشده‌ای که در همسایگی مکانی یا زمانی مشاهدات قرار دارند می‌توانند شامل اطلاعات مفیدی باشند که بازیابی این اطلاعات از دست رفته می‌تواند موجب افزایش دقت تحلیل‌ها شود. در این مقاله تحت فرض گمشدگی غیرقابل چشم‌پوشی، مدل‌بندی توأم فرایند اندازه‌گیری فضایی-زمانی و فرایند گمشدگی با استفاده از تکنیک مدل پارامتر مشترک پیشنهاد شده است، که در آن از طریق یک یا چند متغیر پنهان فضایی-زمانی مشترک، بین دو فرایند ارتباط برقرار می‌شود. استنباط‌ها از طریق رهیافت بیزی و روش تقریب لاپلاس آشیانه‌ای جمع‌بسته و بهره‌گیری از راهکار معادلات دیفرانسیل جزئی صورت گرفته است. سپس مطالعات شبیه‌سازی برای بررسی عملکرد مدل‌بندی توأم انجام شده است.

واژه‌های کلیدی: داده‌های فضایی-زمانی، داده‌های گمشده، مدل‌بندی توأم، مدل اتورگرسیو پویا، تقریب بیزی INLA.
کد موضوع‌بندی ریاضی (۲۰۱۰): 62F15, 91D25, 91B72

۱ مقدمه

عموماً داده‌های فضایی-زمانی مانند داده‌های ثبت میزان و زمان بارندگی، دمای هوا یا آلودگی هوا در محیط‌های غیرآزمایشگاهی و خارج از کنترل محقق بدست می‌آیند. در مواردی وجود ابر، ناصاف بودن آسمان یا حتی بارش سنگین باران می‌توانند از جمله دلایل ایجاد گمشدگی باشند. از طرفی خرابی دستگاه‌ها و ابزارهای اندازه‌گیری نیز می‌توانند منجر به تولید مقادیر گمشده در داده‌ها شوند. **رابین (۱۹۷۶)** فرایندهای مختلف گمشدگی را از هم متمایز نمود و سازوکار گمشدگی را به سه دسته گمشدگی کاملاً تصادفی^۱ (MCAR)، گمشدگی تصادفی^۲ (MAR) و گمشدگی غیرتصادفی^۳ (MNAR) دسته‌بندی کرد.

^۱Missing Completely At Random

^۲Missing At Random

^۳Missing Not At Random

^۱ نام و ایمیل ارائه دهنده مقاله: سمیرا زحمتکش، samira.zahmatkesh@modares.ac.ir

رده‌بندی کرد. وقتی فرایند گمشدگی مستقل از داده‌های مشاهده شده و مشاهده نشده باشد، در این صورت MCAR رخ می‌دهد، رده MAR وقتی وقوع می‌یابد که فرایند گمشدگی تنها به مقادیر مشاهده شده بستگی داشته باشد. مواردی که فرایند گمشدگی در هیچ یک از دو دسته MCAR و MAR نباشد گمشدگی MNAR رخ داده است، در این حالت گمشدگی هم به داده‌های مشاهده شده و هم به داده‌های مشاهده نشده بستگی دارد و اصطلاحاً گمشدگی غیرقابل چشم‌پوشی نامیده می‌شود.

از آنجا که از روی داده‌ها نمی‌توان به سازوکار گمشدگی پی برد، در مواردی که نتوان با قاطعیت ادعا کرد که گمشدگی تصادفی رخ داده است لازم است به منظور دستیابی به نتایج معتبر فرایند گمشدگی به همراه فرایند اندازه‌گیری مدل‌بندی و در تحلیل داده‌ها استفاده شود. در این زمینه روش‌های متعددی بر پایه سه چارچوب مختلف از تجزیه مدل توأم فرایند اندازه‌گیری و فرایند گمشدگی ارائه شده است که یکی از آن‌ها مدل پارامتر مشترک^۴ (SPM) است (دنیل و هوگان، ۲۰۰۸). مدل پارامتر مشترک برای مدل‌بندی داده‌های فضایی گمشده توسط زحمتکش و محمدزاده (۲۰۱۸) مورد بررسی قرار گرفته است. در این مقاله مدل پارامتر مشترک را به داده‌های فضایی-زمانی تعمیم خواهیم داد، که در آن با در نظر گرفتن یک متغیر پنهان فضایی-زمانی مشترک برای توصیف ارتباط بین دو فرایند، تلاش می‌شود بخشی از اطلاعات از دست‌رفته را از طریق داده‌های مشاهده شده بازایی نموده و برای تحلیل در داده‌های فضایی-زمانی مورد استفاده قرار دهیم. در این راستا از مدل‌های اتورگرسیو پویای مرتبه یک فضایی-زمانی برای مدل‌بندی فرایند اندازه‌گیری و فرایند گمشدگی استفاده می‌شود. همچنین در مدل‌بندی و استنباط پارامترهای مدل از روش بیزی تقریب لاپلاس آشیانه‌ای جمع بسته^۵ (INLA) که توسط رو و همکاران (۲۰۰۹) ارائه شده است، استفاده خواهد شد. برای این منظور توزیع توأم دو پاسخی (فرایند اندازه‌گیری و فرایند گمشدگی) در استفاده از INLA در نظر گرفته می‌شود. بعلاوه برای کاهش حجم محاسبات ناشی از ماتریس کواریانس چگال از روشی کارا و نوین از مدل‌بندی فضایی با استفاده از معادلات دیفرانسیل جزئی تصادفی^۶ (SPDE) استفاده می‌شود، که در آن میدان تصادفی گاوسی در فضایی مارکوفی تقریب زده می‌شود و در عین حال امکان استفاده از INLA نیز وجود دارد (لینگرن و همکاران، ۲۰۱۱).

ساختار مقاله به این صورت است که در بخش ۲ مدل‌های فضایی-زمانی معرفی می‌شوند در بخش ۳ مدل توأم فرایند اندازه‌گیری فضایی-زمانی و فرایند گمشدگی معرفی می‌شود. در بخش ۴ مروری کوتاه بر روش INLA و راهکار SPDE خواهیم داشت و در بخش ۵ مطالعه شبیه‌سازی برای بررسی عملکرد مدل توأم اجرا خواهد شد و در انتها به بحث و نتیجه‌گیری پرداخته می‌شود.

۲ مدل‌بندی داده‌های فضایی-زمانی در دامنه پیوسته فضایی

میدان تصادفی فضایی-زمانی $\{y(s, t), (s, t) \in D \subset R^2 \times R\}$ را در نظر بگیرید. فرض کنید مشاهدات در n موقعیت فضایی $s_1, \dots, s_n$ و در T نقطه از زمان، $t = 1, \dots, T$ جمع‌آوری شده‌اند و $y(s_i, t)$ تحقیقی از میدان تصادفی فضایی-زمانی در ایستگاه s_i و در زمان t باشد. مدل فضایی-زمانی به صورت

$$y(s_i, t) = \mathbf{z}(s_i, t)\beta + \xi(s_i, t) + \varepsilon(s_i, t) \quad (1.2)$$

را در نظر بگیرید، که در آن $\mathbf{z}(s_i, t) = (z_1(s_i, t), \dots, z_p(s_i, t))$ بردار مشاهدات p متغیر تبیینی در موقعیت s_i و زمان t ، $\beta = (\beta_1, \dots, \beta_p)'$ بردار ضرایب، $\varepsilon(s_i, t)$ خطای اندازه‌گیری از توزیع $N(0, \sigma_\varepsilon^2)$ و $\xi(s_i, t)$ تحقیقی از یک

^۴Shared Parameter Model

^۵Integrated Nested Laplace Approximation

^۶Stochastic Partial Differential Equations

میدان تصادفی پنهان فضایی-زمانی است. به عنوان مثال در بررسی کیفیت هوا، این میدان تصادفی پنهان می تواند سطح واقعی مشاهده نشده از آلودگی هوا باشد. فرض می شود $\xi(s_i, t)$ در طی زمان بر اساس یک مدل پویای اتورگرسیو مرتبه اول، $AR(1)$ ، به صورت

$$\xi(s_i, t) = a\xi(s_i, t-1) + \omega(s_i, t) \quad (2.2)$$

تغییر می کند (کملتی و همکاران، ۲۰۱۱؛ کوچی و همکاران، ۲۰۰۷؛ شاهو، ۲۰۱۱)، که در آن $|a| < 1$ و $\xi(s_i, 1)$ دارای توزیع $N(0, \frac{\sigma_\omega^2}{1-a^2})$ و $\omega(s_i, t)$ یک میدان تصادفی گاوسی مانا با میانگین صفر و تابع کواریانس فضایی-زمانی

$$Cov(\omega(s_i, t), \omega(s_j, t')) = \begin{cases} 0 & t \neq t' \\ \sigma_\omega^2 \rho(h) & t = t' \end{cases} \quad (3.2)$$

است، که در آن تابع همبستگی فضایی $\rho(h)$ تنها از طریق $h = \|s_i - s_j\| \in R$ به موقعیت های فضایی s_i و s_j وابسته است. به این ترتیب فرض می شود که میدان تصادفی $\omega(s, t)$ مانای مرتبه دوم و همسانگرد است (محمدزاده، ۱۳۹۸). تابع همبستگی فضایی $\rho(h)$ عموماً به دلیل انعطاف پذیری، تابع مترن به صورت

$$\rho(h) = \frac{1}{\Gamma(\nu) 2^{\nu-1}} (\kappa h)^\nu K_\nu(\kappa h) \quad (4.2)$$

در نظر گرفته می شود، که در آن K_ν تابع بسل تعدیل یافته نوع دوم و از مرتبه $\nu > 0$ است (آبرامویتز و استگان، ۱۹۷۲)، پارامتر ν بیانگر درجه همواری فرایند و $\kappa > 0$ پارامتر مقیاس مرتبط با دامنه r است. دامنه فاصله ای است که در آن همبستگی فضایی کوچک و نزدیک به صفر می شود و رابطه $r = \frac{\sqrt{\lambda\nu}}{\kappa}$ برای آن تعریف می شود (لینگرن و همکاران، ۲۰۱۱). فرض کنید $\mathbf{y}_t = (y(s_1, t), \dots, y(s_n, t))$ تحقق از فرایند اندازه گیری در زمان t و n موقعیت فضایی باشد، مدل های (۱.۲) و (۲.۲) را می توان به صورت

$$\mathbf{y}_t = \mathbf{z}_t \beta + \xi_t + \varepsilon_t \quad \varepsilon_t \sim N(0, \sigma^2 I_n) \quad (5.2)$$

$$\xi_t = a\xi_{t-1} + \omega_t \quad \omega_t \sim N(0, R = \sigma_\omega^2 \tilde{R}) \quad (6.2)$$

نوشت، که در آن I_n ماتریس همانی از بعد n است و $\mathbf{z}_t = (z(s_1, t)', \dots, z(s_n, t)')'$ و $\xi_t = (\xi(s_1, t), \dots, \xi(s_n, t))'$. ξ_1 یک فرایند $AR(1)$ از توزیع مانای $N(0, \frac{1}{1-a^2} R)$ ، و $\tilde{R}$ ماتریس همبستگی از بعد n با عناصر $\rho(\|s_i - s_j\|)$ است. روابط (۵.۲) و (۶.۲) مدل های سلسله مراتبی را ارائه می دهند که به دلیل انعطاف پذیری آن ها در مدل بندی همزمان اثرات متغیرهای تبیینی و همبستگی های فضایی-زمانی، مکرراً در مطالعات بررسی کیفیت هوا استفاده شده اند.

فرض کنید $\mathbf{y} = (\mathbf{y}_1, \dots, \mathbf{y}_T)$ بردار مشاهدات متغیر پاسخ و $\xi = (\xi_1, \dots, \xi_T)$ تحقق از میدان تصادفی پنهان فضایی-زمانی در n موقعیت فضایی و T لحظه زمانی است. قرار می دهیم $\theta_1 = (a, \sigma_\omega^2, r)$ و $\theta_2 = (\beta, \sigma_\varepsilon^2)$. بنابراین $\theta = (\theta_1, \theta_2)$ بردار پارامترهای مدل است که باید برآورد شود. با در نظر گرفتن توزیع پیشینی مستقل $\pi(\theta)$ برای پارامترهای مدل و با فرض استقلال شرطی $\mathbf{y}$ روی ξ در زمان و اینکه میدان تصادفی پنهان از یک مدل پویای زمانی مارکوفی پیروی می کند توزیع پسینی توأم پارامترهای مدل به صورت

$$\begin{aligned} \pi(\theta, \xi | \mathbf{y}) &\propto \pi(\mathbf{y} | \xi, \theta) \pi(\xi | \theta_1) \pi(\theta) \\ &\propto \left(\prod_{t=1}^T \pi(\mathbf{y}_t | \xi_t, \theta) \right) \left(\pi(\xi_1 | \theta_1) \prod_{t=2}^T \pi(\xi_t | \xi_{t-1}, \theta_1) \right) \pi(\theta) \end{aligned} \quad (7.2)$$

حاصل می‌شود، که با توجه به (۱.۲) و (۲.۲) داریم

$$\begin{aligned} \pi(\theta, \xi | \mathbf{y}) &\propto (\sigma_\varepsilon^2)^{-\frac{nT}{2}} \exp\left(-\frac{1}{2\sigma_\varepsilon^2} \sum_{t=1}^T (\mathbf{y}_t - \mathbf{z}_t\beta - \xi_t)'(\mathbf{y}_t - \mathbf{z}_t\beta - \xi_t)\right) \\ &\times \left(\frac{\sigma_\omega^2}{1-a}\right)^{-\frac{n}{2}} |\tilde{R}|^{-\frac{1}{2}} \exp\left(-\frac{1-a}{2\sigma_\omega^2} \xi_1' \tilde{R}^{-1} \xi_1\right) \\ &\times (\sigma_\omega^2)^{-\frac{n(T-1)}{2}} |\tilde{R}|^{-\frac{(T-1)}{2}} \exp\left(-\frac{1}{2\sigma_\omega^2} \sum_{t=2}^T (\xi_t - a\xi_{t-1})' \tilde{R}^{-1} (\xi_t - a\xi_{t-1})\right) \\ &\times \prod_{i=1}^{\dim(\theta)} \pi(\theta_i) \end{aligned} \quad (۸.۲)$$

که در آن $|\tilde{R}|$ دترمینان ماتریس همبستگی تنک $\tilde{R}$ است.

۳ مدل‌بندی توأم بیزی فرایندهای گمشدگی و اندازه‌گیری تحت فرض MNAR

در داده‌های فضایی-زمانی این گمان می‌تواند وجود داشته باشد که وجود مقادیر گمشده تحت تأثیر عواملی غیرتصادفی باشند، بنابراین تحت فرض MNAR، یعنی با فرض غیرقابل چشم‌پوشی بودن گمشدگی، می‌توان مدلی توأم از فرایند اندازه‌گیری (متغیر پاسخ) و فرایند گمشدگی ساخت. در اینجا SPM برای حالت فضایی-زمانی بسط داده می‌شود. فرض کنید فرایند گمشدگی به صورت $\mathbf{m} = (\mathbf{m}_1, \dots, \mathbf{m}_T)$ است، که در آن $\mathbf{m}_t = (m(s_1, t), \dots, m(s_d, t))$ ، $t = 1, \dots, T$ از طریق تابع نشانگر

$$m(s_i, t) = \begin{cases} 1 & y(s_i, t) \text{ گمشده است} \\ 0 & y(s_i, t) \text{ مشاهده شده است} \end{cases}$$

تعریف می‌شود، که دارای توزیع برنولی $Ber(p)$ با احتمال گمشدگی p است. در این صورت احتمال گمشدگی برای $y(s_i, t)$ به صورت $P(m(s_i, t) = 1 | \varphi) = p(s_i, t)$ است، که در آن φ بردار پارامترهای مدل فرایند گمشدگی است. احتمال گمشدگی $p(s_i, t)$ از طریق یک رگرسیون لوژیستیک به صورت

$$\text{logit}(p(s_i, t)) = \alpha_0 + \alpha_1 \xi(s_i, t) + \mathbf{H}(s_i, t) \gamma \quad (۱.۳)$$

مدل‌بندی می‌شود، که در آن α_0 عرض از مبدأ و بیانگر نسبت ثابتی از گمشدگی است، $\xi(s_i, t)$ تحقق از یک میدان تصادفی پنهان فضایی-زمانی در موقعیت فضایی s_i و زمان t است که در مدل متغیر پاسخ نیز حضور دارد و در واقع بین $\mathbf{y}$ و $\mathbf{m}$ ارتباط برقرار می‌کند و از طریق ضریب α_1 شدت و چگونگی این ارتباط توصیف می‌شود. این میدان پنهان فضایی-زمانی مشابه (۲.۲) در طی زمان بر اساس اتورگرسیون پویای مرتبه اول تغییر می‌کند. بردار q بعدی $\mathbf{H}(s_i, t) = (H_1(s_i, t), \dots, H_q(s_i, t))$ متغیرهای تبیینی در ارتباط با احتمال گمشدگی است و $\gamma = (\gamma_1, \dots, \gamma_q)$ بردار ضرایب مربوطه است. به این ترتیب $\varphi = (\alpha_0, \alpha_1, \gamma)$. اکنون با وارد کردن مدل گمشدگی در رابطه (۷.۲) به جای توزیع $\pi(\mathbf{y} | \xi, \theta)$ توزیع توأم $\pi(\mathbf{y}, \mathbf{m} | \xi, \varphi, \theta)$ جایگزین خواهد شد. در مدل‌بندی توأم فرایند اندازه‌گیری و فرایند گمشدگی، ξ به عنوان میدان تصادفی پنهان فضایی-زمانی مشترک در نظر گرفته می‌شود و با فرض استقلال شرطی آن دو روی ξ توزیع توأم آن‌ها بر اساس SPM به صورت

$$\pi(\mathbf{y}, \mathbf{m} | \xi, \varphi, \theta) = \pi(\mathbf{y} | \xi, \theta) \pi(\mathbf{m} | \xi, \varphi) \quad (۲.۳)$$

تجزیه می‌شود. تابع درستنمایی داده‌های کامل به صورت

$$\begin{aligned} L(\xi, \varphi, \theta) &= \prod_{t=1}^T \pi(\mathbf{y}_t, \mathbf{m}_t | \xi_t, \varphi, \theta) \\ &= \prod_{t=1}^T \pi(\mathbf{y}_t | \xi_t, \theta) \prod_{t=1}^T \pi(\mathbf{m}_t | \xi_t, \varphi) \end{aligned}$$

خواهد بود. با در نظر گرفتن $\eta = (\varphi, \theta)$ به عنوان بردار پارامترهای مدل توأم، رابطه (۸.۲) به سادگی بازنویسی می‌شود که فرم بسته‌ای ندارد. برای تولید نمونه از توزیع‌های پسینی حاشیه‌ای از روش تقریب بیزی INLA استفاده می‌شود که در بخش بعد معرفی خواهد شد.

۴ تحلیل بیزی با استفاده از SPDE و INLA

مدل سلسه‌مراتبی تعریف شده در (۱.۲) و (۲.۲) عضوی از رده مدل‌های گاوسی پنهان است که از طریق INLA قابل برآورد است که روشی جایگزین برای روش‌های MCMC به منظور دستیابی سریع‌تر به تقریبی از توزیع‌های حاشیه‌ای پسینی برای متغیرهای پنهان و پارامترها است. بر اساس رو و همکاران (۲۰۰۹)، به میدان پنهان مدل توزیع GMRF تخصیص می‌دهیم. یک GMRF یک میدان تصادفی است که وابستگی فضایی داده‌ها را در دامنه‌های فضایی مختلف مدل‌بندی می‌کند (رو و هلد، ۲۰۰۵). فرض کنید $\mathbf{x} = (x_1, \dots, x_n)$ از توزیع $\mathbf{x} \sim N(\mu, \mathbf{Q}^{-1})$ یک GMRF با بعد n باشد، که در آن $\mathbf{Q}$ ماتریس معین مثبت و متقارن دقت، یعنی وارون ماتریس کواریانس است. به این ترتیب تابع چگالی $\mathbf{x}$ به صورت

$$\pi(\mathbf{x}) = (2\pi)^{(-\frac{n}{2})} |\mathbf{Q}|^{\frac{1}{2}} \exp\left(-\frac{1}{2}(\mathbf{x} - \mu)' \mathbf{Q}(\mathbf{x} - \mu)\right)$$

است، که در آن $\mathbf{Q}$ با توجه به خاصیت مارکفی $\mathbf{x}$ ، از طریق ساختار همسایگی مشاهدات تعیین می‌شود. در واقع مزیت استنباط بر اساس GMRF ویژگی محاسباتی خوب و امکان بهبود استنباط بیزی از طریق INLA است (رو و همکاران، ۲۰۰۹). یکی از روش‌های تبدیل GRF به GMRF استفاده از راهکار SPDE است که هدف آن پیدا کردن GMRF با یک همسایگی و ماتریس دقت تنک $\mathbf{Q}$ است به طوری که مشکل n بزرگ (بانرزی و همکاران، ۲۰۰۸) را که هنگام کار با ماتریس کواریانس رخ می‌دهد برطرف می‌نماید. اساساً راهکار SPDE از یک ارایه عنصر متناهی به منظور تعریف GRF به عنوان ترکیب خطی از توابع پایه در یک توری مثلثی از دامنه D استفاده می‌کند. این کار شامل افراز D به مجموعه‌ای از مثلث‌های مجزا است که حداکثر در یک ضلع یا زاویه مشترک هستند. برای تشکیل توری مثلثی ابتدا هر یک از n ایستگاه به عنوان یک رأس مثلث در نظر گرفته می‌شوند، سپس سایر رئوس اضافه شده و مثلث‌ها ساخته می‌شوند. فرض کنید $\mathbf{x}(s) = \{x(s), s \in D \subseteq \mathbb{R}^2\}$ یک میدان تصادفی گاوسی مانای همسانگرد با تابع کواریانس مترن (۴.۲) باشد، با فرض معلوم بودن توری مثلثی، یک SPDE ترکیب خطی از توابع پایه به صورت

$$\mathbf{x}(s) = \sum_{\ell=1}^m \psi_{\ell}(s) w_{\ell} \quad (۱.۴)$$

است، که در آن m تعداد رئوس مثلث‌ها در توری، $\psi_{\ell}(s)$ توابع پایه و w_{ℓ} وزن‌هایی با توزیع گاوسی هستند. آنچه که در راهکار SPDE مورد هدف قرار می‌گیرد ارایه عنصر متناهی (۱.۴) است که ارتباطی بین GRF و GMRF ایجاد می‌کند و از طریق وزن‌های گاوسی w_{ℓ} یک ساختار مارکفی ساخته می‌شود (لینگرن و همکاران، ۲۰۱۱). اکنون در مدل فضایی-زمانی مفروض، در استفاده از راهکار SPDE در هر نقطه از زمان، $t = 1, \dots, T$ میدان مترن w_t از طریق یک GMRF به صورت $\tilde{w}_t \sim N(0, Q_s^{-1})$ ارایه می‌شود، که در آن Q_s ماتریس دقت است. ماتریس Q_s با توجه به (۳.۲) مستقل زمانی است و

بعد آن برابر با تعداد گره‌های توری مثلثی است. بنابراین معادله (۶.۲) برای $t = 1, \dots, T$ به صورت

$$\xi_t = a\xi_{t-1} + \tilde{\omega}_t \quad \tilde{\omega}_t \sim N(0, Q_s^{-1}) \quad (2.4)$$

نوشته می‌شود، که در آن $\xi_1 \sim N(0, \frac{Q_s^{-1}}{1-a^2})$. به این ترتیب توزیع توأم GMRF nT بعدی $\xi = (\xi'_1, \dots, \xi'_T)$ به صورت $\xi \sim N(0, Q^{-1})$ است، که در آن $Q = Q_T \otimes Q_S$ و Q_T ماتریس دقت T بعدی فرایند اتورگرسیو زمانی مرتبه ۱ در (۳.۲) است. بعلاوه معادله (۵.۲) به صورت

$$\mathbf{y}_t = \mathbf{z}_t\beta + B\xi_t + \varepsilon_t \quad \varepsilon_t \sim N(0, \sigma_\varepsilon^2 I_n) \quad (3.4)$$

نوشته می‌شود که در آن ماتریس $n \times m$ بعدی B مقدار ξ_t را برای مشاهده y_t انتخاب می‌کند و یک ماتریس تنک است به طوری که

$$y(s_i, t) = z(s_i, t)\beta + \sum_{j=1}^m B_{ij}\xi_t + \varepsilon_t(s_i, t) \quad (4.4)$$

که در آن $B_{ij} = 1$ اگر رأس j ام در مکان s_i باشد و در غیر اینصورت $B_{ij} = 0$. فرض کنید $\mathbf{x} = \{\xi, \beta\}$ بردار پارامترهای مدل با پیشین‌های مستقل هستند و β توزیع پیشین گاوسی مبهم با ماتریس دقت معلوم و ξ دارای توزیع GMRF است. بنابراین چگالی $\pi(\mathbf{x}|\theta)$ گاوسی با میانگین صفر و ماتریس دقت $Q(\theta_1)$ با ابرپارامتر $\theta_1 = (\sigma_\omega^2, a, \kappa)$ است. اگر $\theta_2 = \sigma_\varepsilon^2$ و $\theta = (\theta_1, \theta_2)$ ، آنگاه مشاهدات $\mathbf{y} = \{\mathbf{y}_t\}$ دارای توزیع نرمال و مستقل شرطی روی $\mathbf{x}$ و θ هستند. در این صورت توزیع پسینی توأم به صورت

$$\pi(\mathbf{x}, \theta|\mathbf{y}) = \pi(\theta)\pi(\mathbf{x}|\theta) \prod_{t=1}^T \pi(\mathbf{y}_t|\mathbf{x}, \theta) \quad (5.4)$$

است، که در آن $\pi(\mathbf{y}_t|\mathbf{x}, \theta)$ توزیع شرطی مشاهدات در زمان t است و بنابر (۳.۴) توزیع نرمال $N(\mathbf{z}_t\beta + B\xi_t, \sigma_\varepsilon^2 I_n)$ است. در اینجا لازم است توزیع‌های پسینی حاشیه‌ای میدان پنهان و ابرپارامترها به صورت

$$\pi(x_i|\mathbf{y}) = \int \pi(x_i|\theta, \mathbf{y})\pi(\theta|\mathbf{y})d\theta, \quad i = (1, \dots, T+p)$$

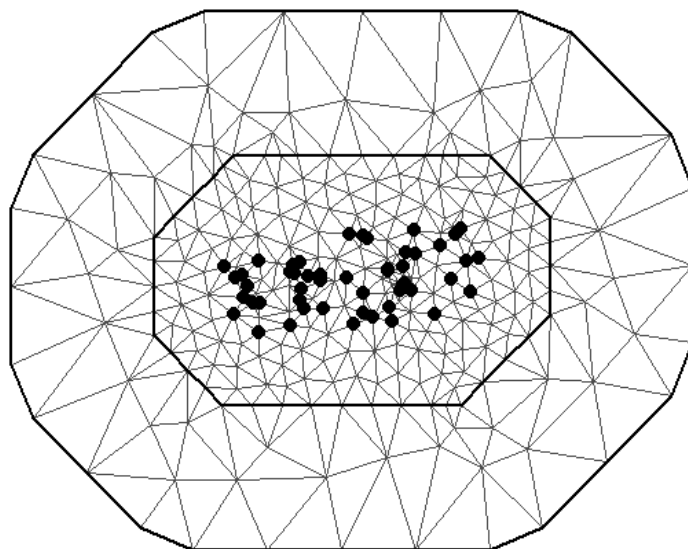
$$\pi(\theta_j|\mathbf{y}) = \int \pi(\theta|\mathbf{y})d\theta_{-j}, \quad j = 1, \dots, 4$$

محاسبه شوند. تقریب INLA از روش‌های تقریب گاوسی و لاپلاس و انتگرال‌های عددی برای محاسبه توزیع‌های پسینی حاشیه‌ای استفاده می‌کند (رو و همکاران، ۲۰۰۹). در مدل‌بندی توأم فرایند گمشدگی و فرایند اندازه‌گیری فضایی-زمانی، از آن‌جا که میدان تصادفی پنهان فضایی-زمانی ξ به عنوان عامل مشترک در مدل (۱.۳) در نظر گرفته شده است، راهکار SPDE که منجر به مدل‌های سلسله مراتبی (۲.۴) و (۳.۴) شد را به مدل (۱.۳) تعمیم داده و استنباط بیزی را اجرا خواهیم نمود.

۵ شبیه سازی مدل توأم فرایندهای گمشدگی و اندازه‌گیری

در این بخش از طریق شبیه‌سازی داده‌های فضایی-زمانی با فرض وجود MNAR در داده‌ها، عملکرد مدل توأم فرایند گمشدگی و فرایند اندازه‌گیری فضایی-زمانی را بررسی نموده و دقت دو مدل توأم و MAR مورد ارزیابی و مقایسه قرار می‌گیرد.

برای بررسی عملکرد مدل توأم دو فرایند گمشدگی و اندازه‌گیری ابتدا $n = 50$ موقعیت فضایی به عنوان ایستگاه‌های اندازه‌گیری به طور نامنظم در دامنه فضایی مورد نظر تولید شده است. موقعیت قرارگیری این نقاط در شبکه مثالی ساخته شده برای راهکار SPDE در شکل ۱ نمایش داده شده است. فرایند تولید در n موقعیت فضایی و $T = 5$ زمان انجام شده



شکل ۱: توری مثالی ساخته شده برای موقعیت‌های فضایی مورد استفاده در راهکار SPDE

است، که در نهایت مجموعه داده شامل $N = n \times T$ مشاهده خواهد بود. مشاهدات فرایند اندازه‌گیری (متغیر پاسخ)، در موقعیت فضایی s_i و زمان t ، بر اساس مدل

$$\begin{aligned} y(s_i, t) &= \beta_0 + \xi(s_i, t) + \varepsilon(s_i, t) \\ \xi(s_i, t) &= a\xi(s_i, t-1) - \omega(s_i, t) \end{aligned} \quad (1.5)$$

تولید شده‌اند، که در آن $\varepsilon \sim N(0, \sigma_\varepsilon^2)$ و $\omega(s_i, t)$ یک میدان تصادفی گاوسی مانا با تابع کواریانس مترن است. همچنین فرایند گمشدگی $m(s_i, t)$ از توزیع برنولی با احتمال گمشدگی $\pi(s_i, t)$ تولید شده است که در آن احتمال گمشدگی واحد i در زمان t به صورت

$$\text{logit}(\pi(s_i, t)) = \alpha_0 + \alpha_1 \xi(s_i, t)$$

مدل‌بندی شده است، ξ به عنوان عامل مشترک با مدل پاسخ بر اساس مدل پارامتر مشترک (SPM) در نظر گرفته می‌شود. به این ترتیب $\theta = (\beta_0, \sigma_\varepsilon, r, \sigma, a, \alpha_0, \alpha_1)$ بردار پارامترهای مدل است که باید برآورد شوند. اجرا در $m = 200$ تکرار انجام شده است و در نهایت از نتایج برآورد پارامترها میانگین گرفته شده است. با توجه به اینکه پارامتر α_1 شدت وابستگی فرایند اندازه‌گیری و فرایند پاسخ را در مدل‌بندی توأم بیان می‌کند، به منظور بررسی رفتار عملکرد مدل دو مقدار $\alpha_1 = 0/2, 0/8$ در نظر گرفته شده است، همچنین به دلیل تأثیری که احتمالاً میدان تصادفی پنهان بر روی الگوی گمشدگی دارد، برای انحراف استاندارد حاشیه‌ای میدان پنهان فضایی دو مقدار $\sigma = 4, 8$ در نظر گرفته شده است. شبیه‌سازی برای دو مجموعه از مقادیر اولیه پارامترهای مدل اجرا شده است. بردار پارامترها در اجرای اول و دوم به ترتیب برابر با

$$\begin{aligned} \theta_1 &= (\beta_0, \sigma_\varepsilon, r, \sigma, a, \alpha_0, \alpha_1) = (-2, 1, 25, 4, 0/6, 0/5, 0/2) \\ \theta_2 &= (\beta_0, \sigma_\varepsilon, r, \sigma, a, \alpha_0, \alpha_1) = (-2, 1, 25, 8, 0/8, 0/5, 0/8) \end{aligned}$$

قرار داده شده‌اند. همان‌طور که ملاحظه می‌شود در اجرای دوم مقادیر اولیه برای پارامترهای α_1 و σ نسبت به اجرای اول افزایش یافته است. لازم به ذکر است بعد از بررسی نتایج در حضور پارامتر α_0 و مشاهده عدم معنی‌داری آن، از مدل کنار گذاشته شده است. برای ارزیابی عملکرد مدل توأم، مدل دیگری تحت عنوان مدل MAR در نظر گرفته شده است. در این مدل فرض می‌شود مکانیزم گمشدگی تصادفی است و با استفاده از یکی از روش‌های رایج، با جانهی مقادیر گمشده، داده‌ها بازسازی شده و با صرف‌نظر کردن از مدل فرایند گمشدگی، مدل اتورگرسیو پویای مرتبه ۱ (۱.۵) به داده‌ها برازش داده شده است. به عنوان نمونه، برآورد پارامترهای دو مدل توأم و MAR برای مقادیر اولیه θ_2 در جدول ۱ گزارش شده است. همان‌طور که ملاحظه می‌شود انحراف استاندارد و میزان اربیی در برآورد پارامترها بر اساس مدل توأم نسبت به مدل MAR کمتر است.

جدول ۱: برآوردهای پسینی پارامترها برای دو مدل توأم و MAR

پارامتر	مقدار اولیه	مدل							
		توأم				MAR			
		برآورد	انحراف استاندارد	۰/۰۲۵	۰/۹۷۵	برآورد	انحراف استاندارد	۰/۰۲۵	۰/۹۷۵
β_0	-۲	-۲/۰۹۹	۰/۸۱۰	-۴/۴۴۱	-۱/۲۵۷	-۱۲/۲۰۱	۱/۶۳۷	-۱۵/۴۰۸	-۸/۹۲۵
σ_e^2	۱	۱/۲۰۲	۱/۶e-۰۷	۱/۱۸۲	۱/۲۴۱	۰/۳۰۲۱	۰/۰۰۵	۰/۰۰۱	۰/۵۱۷
r	۲۵	۲۵/۲۳۲	۱/۳۹۰	۲۲/۵۴۱	۲۷/۹۸۸	۲۴/۰۳۷	۳/۹۰۵	۱۷/۰۷۶	۳۲/۴۰۴
σ	۸	۷/۳۶۵	۰/۶۲۸	۶/۷۲۷	۹/۶۶۴	۱۰/۰۱۴	۱/۰۴۴	۸/۲۸۹	۱۲/۳۸۹
a	۰/۸	۰/۸۱۳	۰/۰۰۱	۰/۷۹۱	۰/۸۷۸	۰/۶۹۱	۰/۰۷۰	۰/۵۱۶	۰/۸۲۲
α_1	۰/۸	۰/۸۲۸	۰/۰۰۲	۰/۴۷۸	۴/۲۷۲	-	-	-	-

برای مقایسه عملکرد دو مدل توأم و MAR در دو اجرای متفاوت درصد پوشش دهی بازه باورمندی در m تکرار برای پارامترهای مدل بدست آمده است. در اجرای ۱ هر دو مدل عملکرد یکسانی دارند، اما در اجرای ۲ که دو پارامتر α_1 و σ افزایش داده شده‌اند میزان پوشش دهی بازه باورمندی در مدل MAR برای دو پارامتر β_0 و a به ترتیب ۴۹ و ۹۵ درصد بدست آمده است. این در حالی است که برای مدل توأم این معیار برای تمامی پارامترها ۱۰۰ درصد بدست آمده است. برای ارزیابی عملکرد پیشگویی این دو مدل، زمان $t = 2$ انتخاب شده است و با در نظر گرفتن $n_p = 10$ موقعیت فضایی به عنوان موقعیت‌های هدف برای پیشگویی، مشاهدات مربوط به این موقعیت‌ها به عنوان مجموعه آزمایشی از مطالعه کنار گذاشته شده و پیشگویی برای دو مجموعه از مقادیر اولیه انجام شده است. سپس معیار جذر میانگین توان‌های دوم خطای پیشگویی‌ها^۷ (RMSE) محاسبه شده است. در اجرای ۱، میزان RMSE برای مدل توأم و مدل MAR به ترتیب برابر با ۴/۹۳۳ و ۴/۳۵۷ بدست آمده و در اجرای ۲، این معیار برای مدل توأم و MAR به ترتیب برابر با ۵/۹۶۶ و ۹/۷۳۰ حاصل شده است. همان‌طور که ملاحظه می‌شود در اجرای ۱، دقت پیشگویی برای دو مدل تقریباً نزدیک به هم و حتی مدل MAR عملکرد نسبتاً بهتری داشته است، ولی برتری عملکرد مدل توأم در اجرای ۲ آشکار می‌شود که در آن پارامتر همبستگی دو فرایند اندازه‌گیری و گمشدگی بزرگتر و تغییرپذیری میدان تصادفی پنهان فضایی بیشتر است.

بحث و نتیجه‌گیری

در این مقاله، تحت فرض گمشدگی MNAR به مدل‌بندی داده‌های فضایی-زمانی پرداخته شد. با توجه به اینکه در برخی مطالعات فضایی-زمانی گمشدگی می‌تواند متأثر از عوامل پنهان فضایی-زمانی باشد، نمی‌توان به سادگی تحت فرض

⁷Root Mean Square Error

MAR به نتایج قابل اعتمادی دست پیدا کرد. نتایج مطالعات شبیه‌سازی در این مقاله نشان می‌دهد که در چنین مواردی می‌توان با بکارگیری تکنیک‌هایی نظیر مدل پارامتر مشترک، دو فرایند گمشدگی و اندازه‌گیری را توأمآً مدل‌بندی کرد و ارتباط عوامل پنهان فضایی-زمانی و فرایند گمشدگی را کشف و بررسی کرد و به نتایج معتبر و قابل اعتمادتری دست یافت.

مراجع

- محمدزاده، م.، (۱۳۹۸)، *آمار فضایی و کاربردهای آن*، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران، چاپ سوم.
- Abramowitz, M. and Stegun, I. (1972), *Handbook of Mathematical Functions*, Courier Dover Publications.
- Banerjee, S., Gelfand, A. E., Finley, A. O. and Sang, H., (2008), Gaussian Predictive Process Models for Large Spatial Datasets, *Journal of the Royal Statistical Society*, **70**, 825-848.
- Cameletti, M., Ignaccolo, R. and Bande, S., (2011), Comparing Spatio-temporal Models for Particulate Matter in Piemonte, *Environmetrics*, DOI: 10.1002/env.1139.
- Cocchi, D., Greco, F. and Trivisano, C., (2007), Hierarchical space-time Modelling of PM10 Pollution, *Atmospheric environment*, **41**, 532-542.
- Daniels, M. J. and Hogan, J. W. (2008), *Missing Data In Longitudinal Studies Strategies for Bayesian Modeling and Sensitivity Analysis*, Chapman and Hall.
- Lindgren, F., Rue, H., and Lindstr, O., (2011), An Explicit Link Between Gaussian Fields and Gaussian Markov Random Fields: The Stochastic Partial Differential Equation Approach, *Journal of the Royal Statistical Society*, **73**, 423-498.
- Rubin, D. B. (1976), Inference and Missing Data. *Biometrika*, **63**, 581-92.
- Rue, H. and Held, L., (2005), *Gaussian Markov Random Fields: Theory and Applications*, Chapman and Hall.
- Rue, H., Martino, S. and N. Chopin, (2009), Approximate Bayesian Inference for Latent Gaussian Models by Using Integrated Nested Laplace Approximations, *Journal of Statistical Planning and Inference*, **137**, 3177-3199.
- Sahu, S. (2011), *Hierarchical Bayesian Models for Space-time Air Pollution Data*, *Handbook of Statistics*, vol 30, Elsevier Publishers, Holland.
- Zahmatkesh, S. and Mohammadzadeh, M. (2018), Bayesian Prediction of Spatial Data with Non-ignorable Missingness, Submitted.



اثر الگوی جغرافیایی بر رضایت مشتری

فاطمه زنجیران^۱، کیومرث مترجم^۲
^۱گروه آمار، دانشگاه علم و فرهنگ
^۲گروه آمار، دانشگاه تربیت مدرس

چکیده:

در این مطالعه، به معرفی مدلی بر مبنای اقتصادسنجی فضایی پرداخته شده است که یک شرکت را قادر می‌سازد الگوهای منطقه‌ای را در داده‌های رضایت‌مندی با استفاده از رگرسیون وزنی جغرافیایی (GWR) شناسایی کند. برای پیاده‌سازی این مدل، از داده‌های مربوط به رضایت‌مندی مشتریان از کیفیت خودرو در ایالات متحده آمریکا استفاده شد. بر مبنای داده‌ها که براساس کدپستی مشتریان است، مختصات طول و عرض جغرافیایی برای هر کدپستی به دست آورده شده است. سپس، باتوجه به فاصله اقلیدسی بین هر کدپستی از مدل GWR استفاده شد. همبستگی فضایی نیز با استفاده از آزمون موران و گری مورد بررسی قرار گرفت و نتایج حاکی از مثبت بودن همبستگی فضایی برای متغیرهای مورد پژوهش بود. سپس، عملکرد مدل براساس توانایی مدل در پیش‌بینی رضایت‌مندی کلی و پایایی برآورد پارامترهای مدل فضایی مورد ارزیابی قرار گرفت.

واژه‌های کلیدی: آمار فضایی، رضایت‌مندی مشتری، رگرسیون وزن جغرافیایی.
کد موضوع بندی ریاضی (۲۰۱۰): 62P20، 62J99، 62J05.

۱ مقدمه

پراکندگی مشتریان شرکت‌ها در مناطق مختلف جغرافیایی سبب دشواری مدیریت کیفیت خدمات می‌شود؛ زیرا همبستگی فضایی بین مناطق موجب ایجاد تفاوت در رضایت‌مندی مشتریان می‌شود. در نظر گرفتن اثرات فضایی در رضایت داده‌ها می‌تواند کیفیت خدمات را افزایش دهد؛ جایی که رضایت‌مندی از کیفیت خدمات پایین است اما مشتریان به دنبال بهبود کیفیت خدمات دریافتی هستند. معمولاً در مناطقی که کیفیت خدمات بسیار بالا است این مطالبه‌گری در مشتریان کم‌رنگ‌تر

^۱ نام و ایمیل ارائه دهنده مقاله: فاطمه زنجیران، n.zanjiran@gmail.com

است. رضایت مشتری برای موفقیت درازمدت شرکت ضروری است (آنتونی و همکاران، ۱۹۹۵) و (والاریه و همکاران، ۲۰۰۰). با این حال، به ویژه در شرکت‌هایی که خدمات را با استفاده از شبکه‌های پراکنده منطقه‌ای از رسانه‌ها (به عنوان مثال، نمایندگی‌های خودرو، شاخه‌های بانک) ارائه می‌دهند، مدیران با اجرای یک استراتژی رضایت مشتری مواجه می‌شوند. برای چنین رسانه‌های پراکنده جغرافیایی، رضایت کلی و اهمیت قرار داده شده بر کیفیت خدمات، منطقه به منطقه متفاوت است. توانایی شرکت در ارائه خدمات برتر نیز ممکن است از لحاظ جغرافیایی متفاوت باشد. به همین ترتیب، ترکیب تفاوت‌های منطقه‌ای در تصمیم‌گیری‌های رضایت مشتری مهم است (ویکاس و واگنر، ۲۰۰۴). در اصطلاح بازاریابی، رضایت مشتری به عنوان اقدامی در مورد چگونگی تحقق محصول یا خدمات ارائه شده توسط سازمان از انتظارات مشتری توصیف می‌شود. این یکی از مهم‌ترین کلیدها برای اطمینان از موفقیت کسب و کار است؛ زیرا رضایت مشتری تعیین‌کننده رشد بازار سازمان در آینده خواهد بود. میزان رضایت از طریق سطح کیفیت محصول، کیفیت خدمات ارائه شده، مکان خرید محصول یا خدمات و قیمت محصول یا خدمات اندازه‌گیری می‌شود (ابراهیم و همکاران، ۲۰۱۶). کیفیت خدمات به عنوان یکی از مهم‌ترین عوامل تعیین‌کننده موفقیت سازمان‌های خدمات در محیط رقابتی امروزه مورد توجه جدی قرار گرفته است. هرگونه کاهشی در رضایت مشتری به دلیل کیفیت پایین خدمت موجب ایجاد نگرانی‌هایی برای سازمان خدماتی است (فولرتون، ۲۰۰۳). در نتیجه بسیاری از متخصصان بازاریابی معتقدند سازمان‌های خدماتی همواره بایستی انتظارات مشتریان از کیفیت خدمات را تحت نظر داشته باشند (پاراسورامان و همکاران، ۱۹۸۵). تعاریف متعددی از کیفیت خدمات ارائه شده است از جمله گرانروس (۲۰۰۱) کیفیت خدمات را اندازه مغایرت بین ادراک مشتری از خدمت و انتظارات او تعریف کرده است (فسنقری و گودرزی، ۱۳۹۲). اگر میانگین نمره رضایت‌مندی برای یک خدمت در یک منطقه خاص بسیار متفاوت باشد، این امر نشان‌دهنده‌ی اشکال در خدمات ارائه شده است و در نهایت منجر به زیان شرکت‌ها می‌گردد؛ و لذا باید رضایت‌مندی افراد در مناطقی که اختلاف سطح خدمات زیاد است، مورد توجه جدی قرار گیرد (برینت و فرای، ۲۰۱۹).

۲ رگرسیون وزنی جغرافیایی

این مقاله به معرفی مدلی بر مبنای اقتصادسنجی فضایی می‌پردازد که یک شرکت را قادر می‌سازد الگوهای منطقه‌ای را در داده‌های رضایت‌مندی با استفاده از رگرسیون وزنی جغرافیایی ^۱ GWR شناسایی کند. برای پیاده‌سازی این مدل از داده‌های مربوط به رضایت‌مندی مشتریان از کیفیت خودرو در ایالات متحده آمریکا استفاده شده است. این رویکرد بر مبنای یک مجموعه داده از ۱۶۴،۰۸۵ مشتری که ۲۱،۶۳۶ کدپستی پنج رقمی در سراسر ایالات متحده را نشان می‌دهد، اعمال شد. مکان‌هایی مورد بررسی قرار گرفته است که در آن رضایت کلی از کیفیت خدمات پایین است اما مطالبه‌گری در آن پررنگ‌تر است. استفاده از روش GWR در تحلیل داده‌های رضایت مشتری به دلیل تطبیق‌پذیری در پرداختن به داده‌های پراکندگی جغرافیایی یک انتخاب مناسب است. تکنیک GWR در حالتی که امکان جمع‌آوری داده‌ها در همه نواحی فراهم نیست مورد استفاده قرار می‌گیرد؛ زیرا این روش از نظر آماری، اطلاعات را برای مناطق همسایه در طول برآورد قرض می‌گیرد (ویکاس و واگنر، ۲۰۰۴). برونسدون و فوترینگام و چارلتون (۱۹۹۶)، توصیف می‌کنند که چگونه GWR از نظر مفهومی شبیه به رگرسیون هسته ^۲ است (رست، ۱۹۸۸). تفاوت اصلی این است که در رگرسیون هسته، وزن‌گیری بر روی "فضای ویژگی" متغیرهای مستقل انجام می‌شود، در حالی که GWR در فضای جغرافیایی دوبعدی انجام می‌شود. تفاوت مهم دیگر این است که در رگرسیون هسته، متغیر وابسته از طریق یک رابطه منفرد، بسیار انعطاف‌پذیر، غیرپارامتری که برای همه مشاهدات یا مکان‌ها اعمال می‌شود، با پیش‌بینی‌کننده‌ها مرتبط است. در مقابل، GWR رابطه خطی بین پیش‌بینی

¹ Graphically Weighted Regression

² Kernel Regression

کننده‌ها و متغیر وابسته است که پارامترها را تخمین می‌زند و در نقاط مختلف متفاوت است. هدف GWR برازش یک مدل خطی است که متغیر وابسته را با متغیرهای تبیینی پس از در نظر گرفتن همبستگی مشاهدات در مکان‌های همسایه مرتبط می‌کند (ویکاس و واگنر، ۲۰۰۴).

۱.۲ شرح مدل

مدل رگرسیون خطی عمومی را در همه‌ی مکان‌ها در نظر بگیرید: $Y = X\beta + \epsilon$ که در آن Y متغیر وابسته و X پیش‌بینی کننده و ϵ عرض از مبدا است. هدف GWR استفاده از تمام اطلاعات موجود در همه نقاط برای برآورد ضرایب رگرسیون است. GWR فرض می‌کند که ضرایب رگرسیون در نقاط مختلف، متفاوت است. ویکاس و واگنر (۲۰۰۴) تکنیک GWR از وابستگی فضایی در داده‌ها استفاده می‌کند. وابستگی فضایی به این معنی است که داده‌های موجود در مکان‌های نزدیک به کانون، اطلاعات بیشتری در مورد ارتباط بین متغیر مستقل و وابسته دارند. GWR هنگام محاسبه برآوردها برای یک موقعیت کانونی، وزن بیشتری به داده‌های مکان‌های نزدیک‌تر نسبت به داده‌های مکان‌های دورتر می‌دهد. از لحاظ آماری، وابستگی فضایی با استفاده از یک طرح وزنی در یک مدل حداقل مربعات عمومی^۳ GLS اجرا می‌شود، به طوری که مکان‌هایی که نزدیک به کانون هستند، وزن بیشتری در تعیین معادله رگرسیون برای محل کانونی دارند. ماتریس وزن حاوی وزن W_j برای همه نقاطی است که در محاسبات معادله رگرسیون برای محل کانونی استفاده می‌شود. با استفاده از این ماتریس وزن جغرافیایی، می‌توان برآورد حداقل مربعات وزنی را برای هر مکان j به صورت زیر به دست آورد:

$$\hat{\beta}_j = (X'W_jX)^{-1}X'W_jy \quad (1.2)$$

این یک رگرسیون عمومی است که با GLS تخمین زده می‌شود، به طوری که W_j ماتریس قطری $n \times n$ است که شامل $W_{jj'}, j' = 1, \dots, j$ در قطر، توسط یک تابع فروپاشی مبتنی بر فاصله تعریف شده است.

$$w_{jj'} = \exp(-d_{jj'}^2/\theta) \quad (2.2)$$

به طوری که θ پارامتر انقطاع فاصله است، و $d_{jj'}$ فاصله اقلیدسی بین مکان‌های j و j' است. پیاده‌سازی تابع فشرده مبتنی بر فاصله، نیاز به برآورد پهنای باند مطلوب (θ) قبل از برآورد حداقل مربعات وزنی دارد. پارامتر پهنای باند ارتباط هر مشاهده همسایه را برای برآورد پارامترهای سطح منطقه‌ای β تعیین می‌کند. هنگامی که پهنای باند به اندازه کافی بزرگ است، مدل GWR به یک رگرسیون استاندارد در سراسر مناطق برمی‌گردد. مناسب‌ترین مقدار پهنای باند با استفاده از روش حداقل مربعات متقاطع که کلیولند (۱۹۷۹) پیشنهاد می‌کند، تعیین شد. اعتبارسنجی متقابل برای مقدار بهینه θ تعیین شد که به صورت تابع امتیاز زیر است:

$$\sum_{j=1}^n [y_j - \hat{y}_{j'}(\theta)]^2 \quad (3.2)$$

که در آن $\hat{y}_{j'}$ مقدار متناسب با y را با مشاهدات مربوط به محل کانونی j که از فرآیند کالیبراسیون حذف شده است، نشان می‌دهد. مقدار θ که این تابع امتیاز را به حداقل می‌رساند به عنوان پهنای باند برای محاسبه ماتریس وزن استفاده می‌شود. جزئیات مربوط به برآورد GWR با اعتبارسنجی متقابل GLS در کار برونسدون و فوترینگام و چارلتون (۱۹۹۶) یافت می‌شود. مدل GWR همانطور که برونسدون و فوترینگام و چارلتون (۱۹۹۸) پیشنهاد می‌کنند، این مورد را تنها با یک

³Generalized Least Squares

مشاهده در هر محل بررسی می‌کنند. با این حال، در این تجزیه و تحلیل، در هر کدپستی چندین پاسخ را مشاهده می‌شود. به جای جمع آوری داده‌ها در هر کدپستی، مدل در سطح پاسخ‌دهندگان تخمین زده می‌شود. این ارزیابی سطح فردی می‌تواند مشکل ساز باشد، زیرا بر اساس تعداد مشاهدات در آن، بر هر مکان تأکید می‌شود. اگرچه این امر در مورد نمونه‌گیری تصادفی ساده یا تصادفی طبقه‌ای ساده قابل قبول است، اما این احتمال وجود دارد که تخمین‌ها را غیر از این سوگیری کند. بنابراین، مفهوم وزن برابر برای هر مکان، مانند مدل GWR عمومی، با وزن هر مشاهده با معکوس اندازه نمونه در محل آن حفظ شده است. مدل GWR تغییر یافته را می‌توان به صورت زیر بیان کرد:

$$\hat{\beta}_j = (X'w_ifX)^{-1}X'w_ify \quad (4.2)$$

که در آن f یک ماتریس قطری $n \times n$ که شامل معکوس اندازه نمونه در مکان j است که مشاهده منحصر به فرد i به آن تعلق دارد و ماتریس قطری w_i شامل وزن‌های مبتنی بر فاصله بین هر مشاهده فردی و کانونی (i) است. توجه شود که اگر ضرایب رگرسیون β_j در سطح مکان تعریف شده باشد، تخمین‌ها برای همه مشاهدات در همان مکان یکسان است؛ زیرا تمام مشاهدات در همان مکان دارای وزن واحد هستند که ناشی از جمع شدن همه مشاهدات از همان مکان می‌شود (ویکاس و واگنر، ۲۰۰۴).

۲.۲ چه زمانی از GWR باید استفاده کرد؟

تکنیک GWR هنگام استفاده از همبستگی فضایی در متغیرها باید مورد استفاده قرار گیرد. همبستگی مثبت و قوی نشان می‌دهد که مقادیر مناطق همسایه مشابه یکدیگر هستند، در حالی که همبستگی منفی و قوی حاکی از آن است که مقادیر مناطق همسایه با یکدیگر متفاوت هستند. بزرگی و جهت همبستگی فضایی برای یک متغیر را می‌توان با استفاده از دو آزمون آماری موران^۴ و گری^۵ اندازه‌گیری کرد. اگر مقادیر موران بزرگ‌تر از $\frac{-1}{N-1}$ باشد، نشان دهنده وابستگی مثبت و معکوس است. برای آزمون گری مقادیر کمتر از ۱ نشان دهنده همبستگی مثبت و مقادیر بیشتر از ۱ نشان دهنده همبستگی منفی فضایی است. برای هر دو آمار، آزمون‌های آماری معنی‌دار می‌تواند برای شناسایی همبستگی فضایی انجام شود. بر اساس این آزمون‌های معنی‌داری آماری، می‌توان تصمیم به ادامه GWR گرفت (کلیف و ارد، ۱۹۸۱، ۱۹۷۳).

۳ داده‌ها

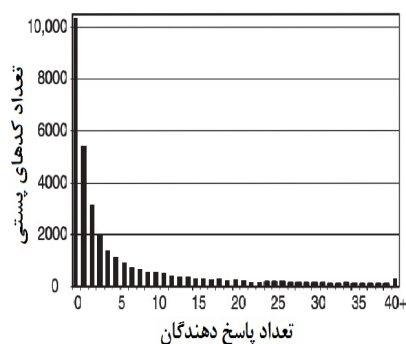
در این قسمت برای پیاده‌سازی مدل معرفی شده از ۱۶۴،۰۸۵ داده مربوط به مشتریانی که در نظرسنجی رضایت‌مندی از کیفیت خودرو شرکت کرده‌اند، استفاده می‌شود. از این مجموعه داده‌ها، نمونه‌ی نهایی به نحوی انتخاب می‌شود که هر کدپستی حداقل شامل ۵ مشتری باشد؛ پس از این شرط، تعداد داده‌ها به ۱۳۲،۰۸۵ مورد کاهش می‌یابد که در مجموع در برگرفته‌ی ۲۱،۶۳۶ کدپستی هستند. در جدول ۱ آمار توصیفی مربوط به نمونه نشان داده شده است. از ۳۱،۹۵۶ کدپستی پنج رقمی در ایالات متحده، تنها ۲۱،۶۳۶ یا ۶۷/۷٪ داده وجود دارد. به عبارت دیگر، ۳۲/۳٪ از کدپستی هیچ داده‌ای ندارند. در شکل ۱ تعداد پاسخ‌دهندگان از هر کدپستی برای کسی که اطلاعات آن در دسترس است، نشان داده شده است. به طور خلاصه، بخش بزرگی از (۳۲/۳٪) کدپستی بدون داده است؛ بیش از نیمی از کدهای پستی باقی مانده، دارای پنج عدد داده یا کمتر هستند. تجزیه و تحلیل سطح کدپستی به طور سیستماتیک تقریباً یک سوم از کدهای پستی را حذف می‌کند. در میان کدهای باقی مانده، برآوردهای سطح کدپستی ممکن است برای کسانی که دارای پنج داده یا کمتر باشند، قابل اعتماد نباشند. مهم‌تر از همه، چنین تحلیلی نمی‌تواند مزایای همبستگی فضایی را در فرآیند برآورد به کار گیرد. به همین

⁴Moran's I

⁵Geary's C

جدول ۱: شرح نمونه

درصد	مشخصات نمونه	
۶۴/۲	مردان	جنسیت
۳۵/۸	زنان	
۲۶/۳	دیپیرستان یا کمتر	سطح تحصیلات
۲۷/۳	تحصیلات دانشگاهی	
۴۶/۵	فارق التحصیل دانشگاه	
۷۶/۰	متاهل	وضعیت تاهل
۱۹/۵	مجرد	
۵/۴	سایر	
۹/۴	زیر ۳۰ سال	سن
۶۸/۱	۳۰-۵۹ سال	
۲۲/۵	بالای ۶۰ سال	
۰/۹	۱	رضایت کلی
۰/۵	۲	
۰/۹	۳	
۲/۱	۴	
۲/۸	۵	
۸/۷	۶	
۹/۴	۷	
۲۵/۴	۸	
۱۸/۵	۹	
۳۰/۷	۱۰	



شکل ۱: فراوانی پاسخ دهندگان با کدپستی

دلیل از GWR برای حل این مشکل استفاده شد: GWR به طور آماری از مشاهدات نقاط همسایه در هنگام برآورد ضرایب برای یک موقعیت کانونی استفاده می‌کند. این ویژگی GWR مفید است، زیرا در اغلب موارد موقعیت کانونی دارای تعداد کمی از مشاهدات است. در چنین مواردی، برآوردهای بهتر را به دست می‌آوریم، زیرا مشاهدات از کدهای پستی همسایه اطلاعات اضافی را فراهم می‌کنند. استفاده از کدهای پستی همسایه نیز به شناسایی مناطقی که دارای ضرایب مشابه هستند،

متغیر	Geary's C	Moran's I	جدول ۲: همبستگی فضایی در متغیرها
رضایت کلی	- ۰/۰۰۸	۰/۶۹	
خدمات نمایندگی	- ۰/۰۰۷	۰/۷۶	
کیفیت خودرو	- ۰/۰۰۷	۰/۷۷	

کمک می‌کند و منجر به دید منظمی از الگوهای فضایی داده‌ها می‌شود.

۱.۳ متغیرها

هر مشتری (i) به سوالات زیر با استفاده از مقیاس ده نمره ($1 =$ "بسیار ناراضی"، $10 =$ "بسیار راضی") پاسخ داد: "بر اساس تجربه، رضایت کلی خود را از مالکیت خودرو ($OVERALLSAT_i$)، کیفیت خودرو ($PRODQUAL_i$) و خدمات نمایندگی ($DLRSRV_i$) چگونه ارزیابی می‌کنید؟"

مشتریان نیز برای پاسخگویی، کدپستی پنج رقمی از محل اقامت فعلی خود را نشان می‌دادند. سپس، مختصات طول و عرض جغرافیایی برای مرکز هر کدپستی پنج رقمی تعیین شد. فاصله اقلیدسی بین هر کدپستی در مجموعه داده محاسبه شد. از این فاصله در برآورد مدل GWR استفاده شد. سپس، همبستگی فضایی در متغیرها تخمین زده شد. اندازه‌گیری همبستگی فضایی (موران و گری) در هر ۲۱,۶۳۶ کدپستی برای هر سه متغیر در جدول ۲ نشان داده شده است. موران بزرگ‌تر از $(\frac{-1}{N-1})$ و $(p < 0/05)$ و گری کمتر از ۱ و $(p < 0/05)$ است. این نشان دهنده همبستگی فضایی مثبت برای هر سه متغیر است.

۴ نتایج

با استفاده از مدل GWR رابطه نهایی به صورت زیر نوشته می‌شود؛ که در آن $OVERALLSAT$ رضایت کلی از تجربه مالکیت، $PRODQUAL$ کیفیت عملکرد خودرو و $DLRSRV$ خدمات نمایندگی است:

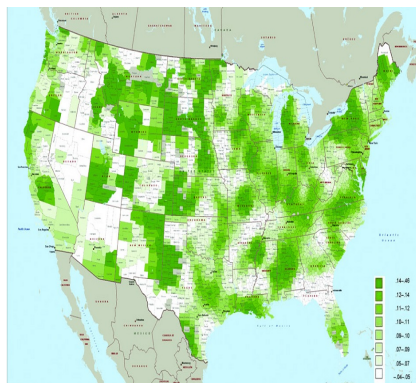
$$OVERALLSAT_i = \beta_{0j} + \beta_{1j}PRODQUAL_i + \beta_{2j}DLRSRV_i + e_i \quad (1.4)$$

که در آن e_i دارای توزیع نرمال است. قابل ذکر است که برای هر فرد i سه پارامتر در هر مکان j تخمین زده شد.

۱.۴ الگوهای منطقه‌ای در اهمیت خدمات نمایندگی

ضریب رگرسیون برای خدمات نمایندگی در شکل ۲ نشان داده شده است. به عنوان مثال در بخش جنوب شرقی کلرادو، اهمیت کیفیت خدمات بالاست (توسط رنگ تیره نشان داده شده)، در حالی که در جنوب غربی کشور، کیفیت خدمات پایین (توسط رنگ روشن‌تر نشان داده شده) است. از مرکز به شمال کلرادو، کم‌ترین اهمیت خدمات مشاهده می‌شود، که توسط سایه سفید نشان داده شده است. علاوه بر این، در بخش شمال شرقی کلرادو، اهمیت رضایت از خدمات به طور یکنواخت پایین است، همانطور که در مجاورت نبراسکا است.

به طور کلی، خدمات نمایندگی در بخش شرقی یا غربی کشور، به‌ویژه در بخش‌هایی از جنوب اورگان و شمال کالیفرنیا، اهمیت بیشتری دارد. سایر مناطق که خدمات نمایندگی در آن مهم است، شمال شرقی نیو مکزیکو تا کلرادو و کانزاس را شامل می‌شود. میدوست، کنتاکی، ایندیانا، و ایلینوی نیز مناطقی هستند که خدمات نمایندگی در آن‌ها از اهمیت بالا یا

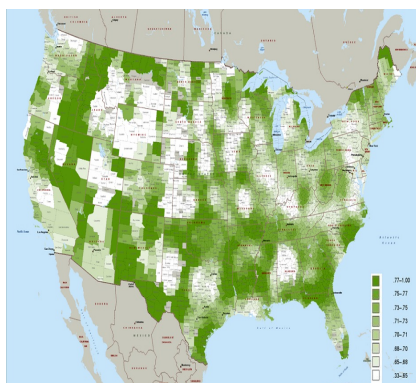


شکل ۲: ضریب اطمینان خدمات نمایندگی

متوسطی برخوردار است. جالب است که یک منطقه بزرگ در غرب آمریکا وجود دارد - از جمله نوادا، آریزونا، اورگان و آیداهو - که خدمات نمایندگی از اهمیت کمتری برخوردار است. علاوه بر این، کلان شهرهای مهم مختلفی در مناطقی قرار گرفته اند که دارای اهمیت متفاوتی هستند. به عنوان مثال ایندیاناپولیس، کلمبوس، اوهایو و فیلادلفیا اهمیت خدمات نمایندگی در آنها زیاد است، در حالی که جکسونویل، فلوریدا، سانفرانسیسکو، سان خوزه و کالیفرنیا، از اهمیت متوسط تا کمی برخوردار است.

۲.۴ الگوهای منطقه‌ای در اهمیت کیفیت خودرو

نتایج مربوط به اهمیت کیفیت وسایل نقلیه در شکل ۳ نشان داده شده است. همان‌طور که انتظار می‌رفت، آنها معکوس نتایج مربوط به خدمات نمایندگی هستند. به عنوان مثال در شمال شرقی تگزاس، اهمیت کیفیت خودرو در تعیین رضایت کلی نسبتاً بالا است. این بخش از تگزاس الگوی اهمیت بالای کیفیت خودرو را با کشورهای همسایه لوئیزیانا و آرکانزاس به اشتراک می‌گذارد. با این حال، در دالاس و آستین، اهمیت نسبی کیفیت خودرو، متوسط رو به پایین است، همان‌طور که در شمال غربی تگزاس این حالت وجود دارد. در مقابل، اهمیت کیفیت خودرو در جنوب تگزاس بسیار زیاد است. الگوی اهمیت کیفیت خودرو برای بازارهای مختلف کلان شهر نیز متفاوت است. در حالی که فیلادلفیا، نیویورک و واشنگتن در مناطق کم اهمیتی قرار دارند، ممفیس و سانفرانسیسکو در مناطقی از اهمیت بالایی برخوردار هستند.



شکل ۳: ضریب اطمینان کیفیت خودرو

جدول ۳: اعتبارسنجی مدل با استفاده از ۲۰٪ نمونه پشتیبان (توانایی پیش‌بینی)

مدل	سطح تجزیه و تحلیل	روش برآورد	دامنه مقدار پیش‌بینی شده به طور متوسط در همه مناطق	درصد مشاهدات واقع در فاصله اطمینان ۹۵٪
۱	۱۳،۸۴۶ کدپستی	GWR	۰/۲۸	۸۷/۰۵
۲	۱۳،۸۴۶ کدپستی	GWR No	۱/۰۷	۴۰/۴۲
۳	شهرستان (تعداد = ۳۱۰۲)	GWR No	۲/۳۹	۷۲/۵۶
۴	وضعیت (تعداد = ۵۰)	GWR No	۳/۸۷	۸۳/۵۸
۵	نمونه از کدپستی (تعداد = ۱۷،۲۲۸)	GWR	۰/۶۲	۷۱/۴۱

۳.۴ عملکرد مدل

عملکرد مدل با استفاده از دو معیار ارزیابی می‌شود؛ توانایی مدل برای پیش‌بینی رضایت کلی و پایایی برآوردهای پارامتر برای رانندگان. برای ارزیابی عملکرد پیش‌بینی شده از نمونه پشتیبان^۶ با استفاده از ۲۰ درصد مشاهدات استفاده شده است. برای ایجاد نمونه پشتیبان، کدهای پستی‌ای انتخاب شد که دارای بیش از پنج مشاهده بود تا بتوان برآوردهای رگرسیونی استاندارد برای هر کدپستی را به دست آورد؛ به همین ترتیب حداقل سه مشاهده برای برآورد مدل‌های معیار نیاز است؛ GWR می‌تواند برآورد پارامتر را حتی در یک منطقه کدپستی که داده‌ای ندارد، بر اساس داده‌ها در مناطق همسایه محاسبه کند. از کدهای پستی که شامل پنج یا بیشتر مشاهده است، به طور تصادفی ۳۲،۰۰۰ مشاهده انتخاب شد؛ کدپستی دارای مشاهدات متعدد در نمونه پشتیبان است. نمونه نهایی حاوی ۳۲،۰۰۰ مشاهده و ۱۳،۸۴۶ کدپستی را نشان می‌دهد.

جدول ۳ GWR را با رگرسیون ساده با استفاده از دو شاخص عملکرد پیش‌بینی شده مقایسه می‌کند. ابتدا، با استفاده از فاصله اطمینان ۹۵٪ برای هر مکان، بهترین گزینه ممکن به عنوان نمونه انتخاب شد. این کار با استفاده از روش پیش‌بینی انجام شد که **نستر (۱۹۹۶)** و **برونسدون و فوترینگام و چارلتون (۱۹۹۸)** پیشنهاد دادند. این سنجش، توانایی مدل را برای پیش‌بینی نمرات رضایت‌مندی بر اساس کدپستی نشان می‌دهد. اندازه‌گیری دوم درصد مشاهدات در نمونه پشتیبان است که برای آنها نمره رضایت واقعی در بازه اطمینان ۹۵٪ توسط مدل پیش‌بینی شده قرار گرفته است. این اندازه‌گیری با نشان دادن درصد دفعات فاصله اطمینان، میزان رضایت واقعی را در یک نمونه پشتیبان، با تأیید تجربی از فاصله اطمینان ۹۵٪ در بین مشتریان ارائه می‌دهد. مدل ۱ از برآورد GWR در سطح کدپستی پنج رقمی استفاده می‌کند. مدل ۲ یک مدل در سطح کدپستی و بدون GWR است که با اجرای رگرسیون‌های جداگانه در هر کدپستی تخمین زده شده است. مدل ۲ حداقل از سه مشاهده در یک کدپستی استفاده می‌کند و برای سایر مشاهدات در کدپستی نمرات رضایت پیش‌بینی شده را با نمونه‌های واقعی مقایسه می‌کند. جدول ۳ نشان می‌دهد که فاصله اطمینان ۹۵٪ برای برآورد مقادیر پیش‌بینی شده با مدل GWR نسبت به مدل بدون GWR کمتر است (۰/۲۸ در مقابل ۱/۰۷). بنابراین، پیش‌بینی‌هایی که با GWR به دست آمده است، نسبت به پیش‌بینی‌های به دست آمده با این فرض که همه مناطق مستقل هستند، بهتر است. ستون سمت چپ جدول ۳ نشان می‌دهد که هنگام استفاده از GWR (مدل ۱)، ۸۷/۰۵٪ از مشاهدات در نمونه پشتیبان دارای نمره رضایت‌مندی هستند که در فاصله اطمینان ۹۵٪ از مقدار پیش‌بینی شده قرار می‌گیرند. بدون GWR (مدل ۲)، فقط ۴۰/۴۲ درصد از مقادیر رضایت کلی در نمونه پشتیبان در فاصله اطمینان ۹۵٪ از مقدار پیش‌بینی شده قرار می‌گیرند. بنابراین، پیش‌بینی‌های بازه‌ای که مدل GWR تولید می‌کند، احتمالاً مقدار واقعی را شامل می‌شود حتی اگر فواصل کوچک‌تر باشند، که نشان‌دهنده عدم قطعیت پایین در مورد پیش‌بینی‌ها است. بازه اطمینان بر اساس رگرسیون مستقل نه تنها گسترده‌تر نبود، بلکه نشان‌دهنده برآورد خطای بالاتر است، همچنین پیش‌بینی رضایت در یک نمونه پشتیبان نیز از دقت کمتری برخوردار بود. برای مقایسه،

^۶Holdout Sample

همچنین مدل‌ها در سطح شهرستان (مدل ۳) و سطح دولت (مدل ۴) تخمین زده شد. همانطور که در جدول ۳ نشان داده شده است، نمره رضایت واقعی در بازه اطمینان ۹۵٪ پیش‌بینی شده برای ۷۲/۵۶٪ از پیش‌بینی‌های انجام شده با برآورد سطح شهرستان و ۸۳/۵۸٪ با مدل سطح دولت می‌باشد. در نگاه اول، به نظر می‌رسد نتایج قابل مقایسه با مواردی است که با برآورد سطح کدپستی به دست آمده است. با این حال، مشاهده می‌کنیم که میانگین دامنه مقادیر پیش‌بینی شده برای هر منطقه برای مدل در سطح شهرستان (۲/۳۹) بسیار گسترده‌تر است و حتی برای مدل سطح دولت (۳/۸۷) حتی بیشتر از مدل تخمین سطح کدپستی است. این بدان معنی است که به احتمال زیاد نمره رضایت واقعی در بازه اطمینان وسیع‌تر قرار می‌گیرد. از آنجا که عدم اطمینان بیشتر در پیش‌بینی‌ها در سطح شهرستان و ایالت وجود دارد، فواصل اطمینان بسیار بیشتر است. بنابراین، برای مدیرانی که می‌خواهند از پیش‌بینی دقیق رضایت کلی برای هر کدپستی مورد نظر اطمینان حاصل کنند، مدل سطح کدپستی با GWR بهترین عملکرد را دارد.

در نهایت، یک نمونه پشتیبان با مجموعه‌ی دیگری از مشاهدات انجام شد. در این حالت به‌طور تصادفی ۳۲،۰۰۰ مشاهده از مجموعه داده‌ها به دست آمد، بدون اینکه محدود به کدپستی‌های شامل ۵ یا بیشتر مشاهده باشد. این نمونه از ۳۲،۰۰۰ نقطه داده شامل ۱۷،۲۲۸ کدپستی است. ۸۰٪ باقی‌مانده از داده‌ها برای تخمین مدل GWR استفاده شد. از مجموعه پارامترهایی که از این مدل به دست آمد، استفاده شد تا پیش‌بینی مشاهدات از نمونه‌های برگزیده انجام شود. نتایج این تحلیل در ردیف پنجم جدول ۳ آمده است. توجه داشته باشید که نتایج این پیش‌بینی با برآوردهای حداقل مربعات معمولی قابل مقایسه نیستند، زیرا برای بسیاری از کدهای پستی، چند نظرسنجی وجود داشت تا مدل حداقل مربعات معمولی برآورد شود. محدوده مقادیر پیش‌بینی شده در حال حاضر بزرگ‌تر از زمانی است که فقط کدپستی با پنج یا بیشتر مشاهدات در دست بود (۰/۶۲ در مقابل ۰/۲۸). این به این دلیل است که کدهای پستی زیادی بود که در منطقه مرکزی هیچ مشاهداتی نداشت و بنابراین باید از طریق مدل GWR پیش‌بینی‌ها از نقاط همسایه تشکیل شود. علاوه بر این، تعداد مشاهداتی که در محدوده $\left(\frac{\sigma}{\sqrt{n}}\right) \pm 1/96$ قرار دارد، ۷۱/۴۱٪ بود، در مقایسه با ۸۷/۰۵٪ برای مدل قبلی. با این حال، این عملکرد پیش‌بینی هنوز هم نسبت به مدل‌های معیار برتر است.

۴.۴ پایداری برآوردهای پارامتر

برای ارزیابی ثبات برآوردها از اهمیت خدمات نمایندگی و کیفیت خودرو، نمونه‌ها در هر کدپستی به صورت تصادفی به دو قسمت تقسیم شد. هنگامی که فقط یک مشاهده برای یک کدپستی در دسترس بود، به صورت تصادفی به یکی از نمونه‌های تقسیم شده اختصاص داده شد، و در نمونه دیگر گم شد. سپس GWR در دو طرف نمونه اعمال شد، و ضرایب کدهای پستی که هیچ داده‌ای نداشتند محاسبه شد. برای ارزیابی ثبات پارامتر، همبستگی از دو طرف و برای داده‌های کامل برای سه مجموعه از کدپستی‌ها محاسبه شد: ۱. آن‌هایی که داده‌هایشان در هر دو نمونه در دسترس بود، ۲. آن‌هایی که داده‌هایشان در یک نمونه گم شده بود، و ۳. آن‌هایی که داده‌هایشان در هر دو نمونه از دست رفته بود (یعنی برآورد پارامتر در هر دو طرف محاسبه شد). تقسیم نیمه ده بار تکرار شد و میانگین و انحراف معیار مربوط به همبستگی در ده تکرار گزارش شد. همبستگی تقسیم نیمه، که در ۳۱،۹۵۶ کدپستی (از جمله کسانی که بدون اطلاعات هستند) محاسبه شده است، ارزیابی ثبات پارامتر برای استفاده از مدل GWR را ارائه می‌دهد. هنگامی که داده‌ها در هر دو طرف تقسیم شده بودند، همبستگی بیشتر از ۰/۶۵ بود. این شواهد قوی در مورد ثبات پارامتر در نظر گرفته می‌شود، زیرا همبستگی بین دو برآورد و یک نمونه بزرگ (۳۱،۹۵۶) محاسبه شده است. وقتی برآورد پارامتر برای یک نمونه تصادفی (یعنی هیچ اطلاعاتی برای کدپستی خاصی در دسترس نیست) محاسبه شود، همبستگی حدود ۰/۴۰ کاهش می‌یابد. این ضعف در تقسیم نیمه متقابل انتظار می‌رفت، زیرا همبستگی‌ها در حال حاضر شامل یک برآورد پارامتر و یک مقدار محاسبه شده است. یک نتیجه قابل توجه این است که همبستگی تقسیم نیمه در زمانی که هر دو نمونه دارای مقادیر محاسبه شده بودند (ستون سمت چپ جدول ۴)

کاهش یافت.

جدول ۴: میانگین و انحراف استاندارد همبستگی تقسیم شده و نیمه (پایداری پارامتر)

همبستگی بین دو نمونه				
پارامتر	رضایت کلی	کدپستی با مشاهده	داده‌های گمشده در یک نمونه	داده‌های گمشده در دو نمونه
عرض از مبدا	۰/۶۱ و ۰/۰۲	۰/۶۶ و ۰/۰۴	۰/۴۱ و ۰/۰۴	۰/۴۱ و ۰/۰۸
خدمات نمایندگی	۰/۶۲ و ۰/۰۴	۰/۶۸ و ۰/۰۶	۰/۴۱ و ۰/۰۵	۰/۴۱ و ۰/۰۷
کیفیت خودرو	۰/۵۸ و ۰/۰۶	۰/۷۰ و ۰/۰۸	۰/۴۱ و ۰/۰۶	۰/۴۶ و ۰/۱۱

بحث و نتیجه‌گیری

در این مقاله براساس روش GWR، نتایج (برای یک شرکت در صنعت خودرو) موارد زیر را نشان می‌دهد: یک تنوع مکانی سیستماتیک در الگوی رضایت کلی و اهمیتی که در رانندگان آن به وجود می‌آید، وجود دارد. اگرچه الگوی خاص تغییرات منطقه‌ای براساس دسته مورد بررسی متفاوت است، اما وجود تغییرات مکانی منظمی باید در تحقیقات بیشتری از داده‌های رضایت گنجانیده شود. اختلافات منطقه‌ای در اهمیت رانندگان و الگوهای کلی رضایت‌مندی یک شرکت را قادر می‌سازد مناطقی که باید در آن خدمات را بهبود بخشند، شناسایی کند. این شرکت می‌تواند در اولویت بندی مناطق برای سرمایه گذاری در بهبود خدمات نمایندگی اقدام کند.

مراجع

فسنقری، ا.، گودرزی، م.، (۱۳۹۲)، بررسی تاثیر ابعاد کیفیت خدمات مدل سروکوال بر رضایت‌مندی مشتریان زن باشگاه‌های زنان، نشریه پژوهش‌های فیزیولوژی و مدیریت در ورزش، ۹، ۲۱-۳۴.

Anthony, J., Zahorik. and Timothy, L., Keiningham. (1995), Return on Quality: Making Service Quality Financially Accountable, *Journal of Marketing*, **59**, 58–70.

Brint, A., Fry, J. (2019), Regional Bias When Benchmarking Services Using Customer Satisfaction Scores, *Total Quality Management and Business Excellence*, 1-15.

Bronsdon, A., Fotheringham, S., Martin, E and Charlton. (1996), Geographically Weighted Regression: A Method for Exploring Spatial Nonstationarity, *Geographical Analysis*, **28**, 281–298.

Bronsdon, A., Fotheringham, S., Martin, E and Charlton. (1998), Geographically Weighted Regression: Modeling Spatial Nonstationarity, *Journal of the Royal Statistical Society D*, **47**, 431–443.

Cliff., Andrew, D. and Ord, J.K. (1973), *Spatial Autocorrelation*, London.

Cliff., Andrew, D. and Ord, J.K. (1981), *Spatial Processes: Models and Applications*, London.

- Cleveland., William, S. (1979), Robust Locally Weighted Regression and Smoothing Scatterplots, *Journal of the American Statistical Association*, **74**, 829–836.
- Fullerton, G. (2003), Women Sport Marketing Business, **124**, 42-48.
- Gronroos, C. (2001), *Service Management and Marketing*, Second Editions, Wiley, 95.
- Iberahim, H., Mohd Taufik, N.K., Mohd Adzmir, A.S. and Saharuddin, H. (2016), Customer Satisfaction on Reliability and Responsiveness of Self Service Technology for Retail Banking Services, *Procedia Economics and Finance* , **37**, 13 – 20.
- Nester., Marks, R. (1996), An Applied Statistician's Creed, *Applied Statistics*, **45 (4)**, 401–410.
- Parasuraman, A., Valarie, A., and Berry, L. (1985), A Conceptual Model of Service Quality and its Implications for Future Research, *Journal of marketing*, **49**, 41-50.
- Rust., Roland, T. (1988), Flexible Regression, *Journal of Marketing Research*, **25**, 10–24.
- Valarie, A., Zeithaml, and Katherine, N., Lemon. (2000), How Customer Lifetime Value Is Reshaping Corporate Strategy, *Driving Customer Equity*, New York: The Free Press.
- Vikas, M., Wagner, A., Kamakura. and Rahul, G. (2004), Geographic Patterns in Customer Service and Satisfaction: An Empirical Investigation, *Journal Of Marketing*, **68**, 48-62.



تحلیل داده‌های فضایی-زمانی بر اساس تجزیه تاکر تانسور کوواریانس

سمیرا سعادت^۱، محسن محمدزاده
گروه آمار، دانشگاه تربیت مدرس

چکیده:

در تحلیل داده‌های فضایی-زمانی، روش متداول برای لحاظ کردن ساختار همبستگی داده‌ها، استفاده از تابع کوواریانس فضایی-زمانی است، که معمولاً نامعلوم است و باید بر اساس مشاهدات برآورد شود. برای این منظور معمولاً محدود کننده‌هایی از قبیل مانایی، همسانگردی و تفکیک‌پذیری میدان تصادفی در نظر گرفته می‌شود که برازش مدل‌های معتبر به تابع کوواریانس فضایی-زمانی را تسهیل می‌کنند. اما در اغلب مسائل کاربردی این فرض‌ها لزوماً محقق نمی‌شوند. در این مقاله، به منظور رهایی از این مشکل و تسریع محاسبه پیشگویی فضایی-زمانی برای یک میدان تصادفی، یک مدل احتمالی بر اساس تجزیه تاکر تانسور کوواریانس فضایی-زمانی ارائه می‌شود، سپس نحوه کاربست روش مطرح شده در پیشگویی فضایی-زمانی داده‌های سرعت باد ایران نشان داده خواهد شد.

واژه‌های کلیدی: داده‌های فضایی-زمانی، تانسور کوواریانس، تجزیه تاکر.
کد موضوع بندی ریاضی (۲۰۱۰): 60Bxx, 60Hxx, 60Gxx.

۱ مقدمه

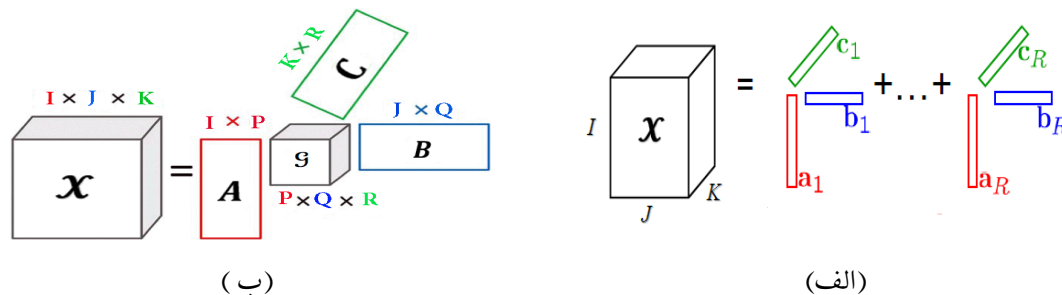
تانسور تعمیمی از مفاهیم اسکالر، بردار و ماتریس است که ساده‌ترین حالت آن می‌تواند دارای یک عضو باشد و تانسور اسکالر نامیده می‌شود. تانسور مرتبه یک، یک بردار، تانسور مرتبه دو، یک ماتریس و تانسورهای مراتب سه و بالاتر تانسورهای مرتبه بالا نامیده می‌شوند. تانسورها از آن جهت مورد استفاده قرار می‌گیرند که باعث ایجاد نظم بین داده‌های یک مسئله و دسته‌بندی اطلاعات آن می‌شوند و بستر ریاضی مناسب و ساده‌ای را برای فرمول‌بندی و حل مسائل متعدد در مباحث گوناگون نظیر مکانیک سیالات و نسبیت عام فراهم می‌کنند. تجزیه تانسور اولین بار توسط **هیچکاک (۱۹۲۷)** معرفی شد، اما کمتر مورد توجه قرار گرفت تا زمانی که **تاکر (۱۹۶۳)** و **کراسکال (۱۹۷۷)** روش‌های مختلف تجزیه تانسور را مطرح کردند و باعث

¹Canonical Polyadic

²Tucker decomposition

^۱ نام و ایمیل ارائه دهنده مقاله: سمیرا سعادت، s.saadati@modares.ac.ir

انجام پژوهش‌های متنوعی بر اساس تجزیه تانسور شد. دو نوع اصلی از تجزیه تانسورها، تجزیه چندتایی کانونی^۱ (CP) و تجزیه تاکر^۲ هستند.



شکل ۱: تجزیه‌های تانسور مرتبه سه، الف- تجزیه CP، ب- تجزیه تاکر

تجزیه CP، یک تانسور را به مجموع تانسورهای با رتبه یک تجزیه می‌کند. برای مثال، یک تانسور مرتبه سه $\mathcal{X} \in \mathbb{R}^{I \times J \times K}$ با استفاده از تجزیه CP به صورت

$$\mathcal{X} \approx \sum_{r=1}^R a_r \circ b_r \circ c_r = \llbracket A, B, C \rrbracket \quad (1.1)$$

تجزیه می‌شود، که در آن R یک عدد صحیح مثبت،

$$a_r \in R^I, b_r \in R^J, c_r \in R^K, \quad r = 1, \dots, R.$$

نماد $\circ$ نشان‌دهنده ضرب خارجی بردارها است. مؤلفه‌ای تانسور (۱.۱) به صورت

$$\mathcal{X}_{ijk} \approx \sum_{r=1}^R a_{ir} b_{jr} c_{kr} \quad i = 1, \dots, I, j = 1, \dots, J, k = 1, \dots, K.$$

نوشته می‌شود، که در شکل ۱- الف نشان داده شده است. تجزیه تاکر نیز هر تانسور مرتبه سه را به یک تانسور هسته ضرب شده در سه ماتریس مانند شکل ۱- ب به صورت

$$\mathcal{X} \approx G \times_1 A \times_2 B \times_3 C = \sum_{p=1}^P \sum_{q=1}^Q \sum_{r=1}^R g_{pqr} a_p \circ b_q \circ c_r = \llbracket G; A, B, C \rrbracket \quad (2.1)$$

تجزیه می‌کند، که در آن $A \in R^{I \times P}$ ، $B \in R^{J \times Q}$ و $C \in R^{K \times R}$ ماتریس‌های فاکتور و $G \in \mathbb{R}^{P \times Q \times R}$ تانسور هسته است و ورودی‌های آن، سطح تعامل بین مؤلفه‌های مختلف را نشان می‌دهد. تجزیه عنصر به عنصر تاکر مرتبه سه در (۲.۱) به صورت

$$x_{ijk} \approx \sum_{p=1}^P \sum_{q=1}^Q \sum_{r=1}^R g_{pqr} a_{ip} \circ b_{jq} \circ c_{kr}, \quad i = 1, \dots, I, j = 1, \dots, J, k = 1, \dots, K,$$

است، که در آن P ، Q و R تعداد مؤلفه‌ها (ستون‌ها) به ترتیب در ماتریس‌های فاکتور A ، B و C هستند.

پیشگویی قابل اعتماد سرعت وزش باد در برخی حوزه‌ها از قبیل ایمنی هوانوردی، ایمنی سازه‌های ساختمانی، انرژی باد و انتقال آلاینده‌ها بسیار مهم است.

اغلب ابزارهای پیشگویی سرعت باد مبتنی بر فیزیک (لینچ، ۲۰۰۸) یا داده‌های آماری (سومن و همکاران، ۲۰۱۰) هستند. گاهی هم ترکیبی از مدل‌های فیزیکی و ابزارهای آماری هستند (سالدو-سانز و همکاران، ۲۰۰۹). عصمتی و محمدزاده ۱۳۹۶ ساختار همبستگی فضایی-زمانی یک میدان تصادفی نامانا و تفکیک‌ناپذیر را با استفاده از تانسور کوواریانس فضایی-زمانی تعیین و پیشگویی فضایی-زمانی مبتنی بر تجزیه CP روی مجموعه داده‌های سرعت باد در کشور ایرلند را انجام دادند. در این مقاله، به منظور رسیدن به نتایج دقیق‌تر به بررسی ساختار همبستگی فضایی-زمانی یک میدان تصادفی نامانا و تفکیک‌ناپذیر (محمدزاده، ۱۳۹۸) با استفاده از تجزیه تانسور کوواریانس فضایی-زمانی مبتنی بر تجزیه تاکر پرداخته خواهد شد و در مطالعه‌ای شبیه‌سازی دقت پیشگویی‌ها مبتنی بر تجزیه CP و تجزیه تاکر مورد ارزیابی و مقایسه قرار می‌گیرد. در انتها نحوه بکارگیری روش ارائه شده برای تحلیل داده‌های سرعت باد ایران نشان داده خواهد شد.

۲ معرفی مدل

فرض کنید سرعت باد در موقعیت فضایی s و لحظه زمانی t با میدان تصادفی $Z(s, t, \theta)$ نشان داده شود، که در آن θ نشان‌دهنده عدم قطعیت موجود در آن است. فرض کنید سرعت باد در موقعیت‌های فضایی $s_i : i = 1, \dots, N_s$ و لحظه‌های زمانی $t_j : j = 1, \dots, N_t$ اندازه‌گیری شده است و مقادیر آن در بردار

$$Z(\theta)^T = \{Z(s_1, t_1, \theta), Z(s_2, t_1, \theta), \dots, Z(s_{N_s}, t_1, \theta), Z(s_1, t_2, \theta), \dots, Z(s_{N_s}, t_{N_t}, \theta)\}^T$$

ارائه شده است. تابع کوواریانس فضایی میدان در لحظه زمانی t به صورت

$$Cov(Z(s_i, t, \theta), Z(s_j, t, \theta)) = E\{[Z(s_i, t, \theta) - \bar{Z}(s_i, t)][Z(s_j, t, \theta) - \bar{Z}(s_j, t)]\} \quad (۱.۲)$$

تعریف می‌شود، که ماتریس کوواریانس آن به صورت تانسور مرتبه سه با اندازه $N_s N_t \times N_s N_t$ است و با C نمایش داده می‌شود. اکنون پیشگویی میدان تصادفی بر اساس تابع کوواریانس (۱.۲) طی شش مرحله ارائه می‌شود.

مرحله ۱- جمع‌آوری و آماده‌سازی داده‌ها: در صورت وجود روندهای فصلی یا دوره‌ای، داده‌ها روند زودوده شوند.

مرحله ۲- برآورد کوواریانس: بر اساس داده‌های روند زودوده ماتریس کوواریانس تجربی یا تانسور محاسبه شود.

مرحله ۳- تجزیه تانسور ماتریس کوواریانس: درایه‌های ماتریس تانسور کوواریانس C به صورت

$$c_{i,j,k} = E\{(Z(s_i, t_k, \theta) - \bar{Z}(t_k, \theta))(Z(s_j, t_k, \theta) - \bar{Z}(t_k, \theta))\} \quad (۲.۲)$$

ایجاد شود.

مرحله ۴- تجزیه‌های تانسور CP و تاکر ماتریس کوواریانس C به صورت

$$\hat{C}_{CP} = [\sum_{r=1}^R \lambda_r \cdot a_r \circ b_r] \circ c, \quad \lambda_r \in R, \quad a, b \in R^{N_s}, \quad c \in R^{N_t} \quad (۳.۲)$$

$$\hat{C}_{Tucker} = [\sum_{p=1}^P \sum_{q=1}^Q \sum_{r=1}^R g_{pqr} \cdot a_p \circ b_q] \circ c, \quad g_{pqr} \in R, \quad a, b \in R^{N_s}, \quad c \in R^{N_t} \quad (۴.۲)$$

محاسبه شوند، که در آن $\hat{C}$ تقریبی از تانسور کوواریانس C است.

مرحله ۵- خطای تقریب $\hat{C}$ به صورت

$$er_C = \frac{\|C - \hat{C}\|_{Fr}}{\|C\|_{Fr}} \quad (۵.۲)$$

محاسبه شود، که در آن $\|.\|_{Fr}$ نرم فروبنیوس است که به صورت

$$\|C\|_{Fr} = \sqrt{\sum_{i=1}^{I_1} \sum_{j=1}^{I_2} \sum_{k=1}^{I_3} c_{i,j,k}^2}$$

تعریف می‌شود (کلدا و بادر، ۲۰۰۹).

مرحله ۶- پیش‌گویی میدان تصادفی مبتنی بر دو تجزیه CP و تاکر به ترتیب به صورت

$$Z_{CP}(\theta) = \bar{Z} + \sum_{r=1}^R \sqrt{\lambda_r} a_r \sqrt{c^T} \eta_r, \quad \lambda_r \in R, \quad a \in R^{N_s}, \quad c \in R^{N_t} \quad (6.2)$$

$$Z_{Tucker}(\theta) = \bar{Z} + \sum_{p=1}^P \sum_{q=1}^Q \sum_{r=1}^R \sqrt{g_{pqr}} a_p \sqrt{c^T} \eta_{pqr}, \quad g_{pqr} \in R, \quad a \in R^{N_s}, \quad c \in R^{N_t} \quad (7.2)$$

محاسبه شود، که در آن η_r ها عضوهای یک مجموعه از متغیرهای تصادفی مستقل با میانگین صفر هستند.

۳ مطالعه شبیه‌سازی و تحلیل داده‌های واقعی

در این بخش طی یک مطالعه شبیه‌سازی عملکرد پیش‌گویی مبتنی بر دو تجزیه CP و تاکر براساس میانگین توان دوم خطای پیش‌گویی مورد ارزیابی و مقایسه عددی قرار می‌گیرد. در این مطالعه ۱۰ موقعیت از بین ۱۰۰ موقعیت جغرافیایی مرز ایران انتخاب شده است. ساختار همبستگی فضایی-زمانی تفکیک‌ناپذیر آن از مدل نایتینگ (۲۰۰۲) به صورت

$$C(h, u) = \left(\frac{\sigma^2}{1 + u^\lambda} \right) \exp \left(\frac{-h^\nu}{(1 + u^\lambda)^{\gamma \delta \nu}} \right), \quad \lambda, \nu \in [0, 2], \gamma \in [0, 1].$$

در نظر گرفته شده است، که در آن h و u به ترتیب تأخیرهای فضایی و زمانی هستند. بر اساس این تابع کوواریانس و موقعیت‌های در نظر گرفته شده یک میدان تصادفی گاوسی تولید می‌شود. هر یک از موقعیت‌ها به ترتیب کنار گذاشته شده و میانگین توان دوم خطای اعتبارسنجی متقابل (CVMSE) پیش‌گویی‌ها مبتنی بر تجزیه‌های CP و تاکر به ازای مقادیر مختلف $\lambda = 0.1, 1, 2$ ، $\nu = 0.1, 1, 2$ و $\sigma^2 = 0.5$ محاسبه شده است. با مقایسه مقادیر میانگین توان دوم خطای اعتبارسنجی متقابل پیش‌گویی‌ها به ازای مقادیر $\nu = 0.1, 1, 2$ همواره پیش‌گویی مبتنی بر تجزیه تاکر نسبت به تجزیه CP از دقت بیش‌تری برخوردار است.

در ادامه، داده‌های متوسط سرعت باد روزانه از سال ۲۰۰۰ تا ۲۰۱۶ در سی و یک ایستگاه بادسنجی کشور که از سایت <https://www.tutiempo.net/clima> قابل دسترسی است، تحلیل آماری شده و با استفاده از مدل ارائه شده اقدام به پیش‌گویی فضایی-زمانی متوسط هفتگی سرعت باد می‌شود. برای این منظور با حذف روند فصلی داده‌ها، تانسور کوواریانس با درایه‌های (۲.۲) بر اساس داده‌ها برآورد شده و به‌صورت‌های (۳.۲) و (۴.۲) تجزیه شده اند. سپس با استفاده از نرم افزار Matlab مقدار $R = 52$ برای رتبه تجزیه CP و رتبه $R = [24, 24, 49]$ برای تجزیه تاکر حاصل شده است. خطای نسبی بین تانسور C و $\hat{C}$ بر طبق (۵.۲) برای تجزیه‌های تاکر و CP به ترتیب ۰/۰۲ و ۰/۱ شده است. همان‌طور که ملاحظه می‌شود، تقریب تانسور C با تجزیه تاکر خطای کمتری نسبت به تقریب حاصل از تجزیه CP دارد. برای پیش‌گویی میدان تصادفی از روابط (۶.۲) و (۷.۲) استفاده شده است، که در آن η_r دارای توزیع نرمال استاندارد منظور گردیده است. نتایج پیش‌گویی متوسط سرعت باد در پنج هفته اول سال ۲۰۱۶ همه ایستگاه‌ها در جدول ۲ ارائه شده است. همان‌طور که ملاحظه می‌شود مقادیر میانگین توان دوم خطای اعتبارسنجی متقابل پیش‌گویی‌های مبتنی بر تجزیه تاکر کوچکتر از مقادیر متناظر حاصل از تجزیه CP هستند، بنابراین پیش‌گویی‌های مبتنی بر تجزیه تاکر از

جدول ۱: میانگین توان دوم خطای اعتبارسنجی متقابل برای پیش‌گویی‌ها براساس تجزیه‌های CP و تاکر

ν	موقعیت	تجزیه CP			تجزیه تاکر		
		$\sigma^2 = 0.5$	$\sigma^2 = 1$	$\sigma^2 = 2$	$\sigma^2 = 0.5$	$\sigma^2 = 1$	$\sigma^2 = 2$
۰/۱	۱	۰/۴۵	۱/۰۵	۱/۲۳	۰/۴۲	۰/۹۶	۱/۱۳
	۲	۰/۶۰	۱/۴۴	۲/۴۴	۰/۳۷	۱/۳۶	۱/۸۰
	۳	۰/۸۷	۱/۱۳	۲/۰۳	۰/۸۴	۰/۹۷	۱/۵۸
	۴	۰/۴۷	۰/۶۴	۲/۹۰	۰/۴۲	۰/۵۳	۱/۵۲
	۵	۰/۵۴	۱/۱۷	۳/۲۱	۰/۳۹	۰/۹۵	۳/۰۷
	۶	۰/۶۹	۱/۲۲	۲/۰۹	۰/۴۷	۰/۸۱	۱/۲۶
	۷	۰/۴۰	۰/۹۰	۰/۹۸	۰/۳۷	۰/۷۱	۰/۹۶
	۸	۰/۶۷	۰/۸۲	۲/۰۵	۰/۴۵	۰/۷۱	۱/۰۲
	۹	۰/۶۲	۱/۷۴	۳/۰۶	۰/۴۴	۱/۱۰	۱/۲۹
	۱۰	۰/۶۹	۰/۸۵	۱/۳۰	۰/۶۷	۰/۷۴	۰/۹۳
۱	۱	۰/۴۶	۱/۱۲	۱/۳۳	۰/۴۴	۰/۹۶	۱/۲۰
	۲	۰/۹۳	۱/۰۷	۱/۹۸	۰/۱۰	۰/۹۸	۱/۴۷
	۳	۰/۶۵	۱/۱۰	۲/۴۷	۰/۵۶	۱/۰۲	۱/۲۸
	۴	۰/۶۲	۰/۸۶	۱/۸۴	۰/۵۸	۰/۷۹	۱/۲۶
	۵	۰/۲۶	۰/۵۹	۱/۱۹	۰/۱۱	۰/۴۹	۰/۹۴
	۶	۰/۵۸	۱/۰۸	۱/۶۱	۰/۳۸	۰/۶۶	۱/۲۸
	۷	۰/۴۶	۱/۱۷	۲/۰۳	۰/۱۶	۰/۹۲	۱/۲۳
	۸	۰/۸۶	۲/۱۷	۲/۵۲	۰/۱۴	۱/۰۱	۱/۸۱
	۹	۰/۵۰	۱/۲۴	۳/۰۳	۰/۴۴	۰/۸۴	۱/۴۸
	۱۰	۰/۸۵	۱/۷۹	۲/۵۳	۰/۷۰	۱/۰۹	۱/۹۳
۲	۱	۰/۹۶	۲/۳۰	۲/۳۶	۰/۷۴	۱/۴۶	۱/۸۵
	۲	۱/۲۰	۱/۳۴	۲/۵۵	۰/۶۶	۱/۱۸	۱/۷۲
	۳	۱/۳۰	۱/۵۵	۱/۸۴	۰/۸۹	۰/۹۷	۱/۶۶
	۴	۱/۷۰	۲/۵۲	۲/۹۱	۰/۴۴	۱/۰۶	۱/۸۰
	۵	۲/۰۵	۲/۶۶	۳/۵۰	۱/۴۲	۱/۵۵	۱/۸۳
	۶	۱/۳۴	۱/۶۷	۳/۳۸	۰/۸۴	۱/۵۵	۲/۳۴
	۷	۱/۰۶	۱/۷۲	۲/۲۵	۰/۷۰	۱/۶۲	۱/۸۴
	۸	۱/۹۱	۲/۱۴	۲/۲۰	۱/۳۳	۱/۴۳	۱/۷۰
	۹	۱/۸۴	۲/۷۴	۳/۳۵	۰/۶۲	۱/۷۰	۲/۶۶
	۱۰	۱/۰۹	۱/۸۰	۲/۶۷	۰/۷۲	۱/۶۱	۲/۰۵

دقت بیش‌تری نسبت به تجزیه CP برخوردار هستند. باید توجه شود که تغییرپذیری زیادی که در سرعت باد پیش‌گویی شده گزارش شده است، ناشی از مدل پیش‌گویی نیست، بلکه به دلیل عدم قطعیت موجود در داده‌های واقعی است که توسط پارامتر θ در مدل (۲.۲) قابل بررسی است. به عبارت دیگر این تغییرپذیری زیاد موجود در داده‌های واقعی است که در پیش‌گویی‌ها نیز منعکس شده است.

جدول ۲: پیش‌گویی متوسط سرعت باد پنج هفته اول سال ۲۰۱۶

ردیف	نام ایستگاه	مقدار واقعی	تجزیه CP		تجزیه تاکر	
			پیش‌گویی	CVMSE	پیش‌گویی	CVMSE
۱	آذربایجان شرقی	۲/۷۷	۲/۱۳	۱/۵۶	۲/۴۲	۰/۴۴
۲	آذربایجان غربی	۲/۱۳	۲/۴۰	۲/۱۳	۲/۲۰	۱/۲۲
۳	اردبیل	۱/۲۷	۱/۴۴	۱/۲۳	۱/۲۹	۱/۱۱
۴	اصفهان	۱/۷۳	۱/۲۱	۱/۱۴	۱/۸۳	۰/۸۵
۵	البرز	۲/۶۲	۲/۵۴	۱/۶۶	۲/۵۸	۱/۶۰
۶	ایلام	۲/۳۷	۲/۱۴	۲/۱۲	۲/۳۰	۰/۷۵
۷	بوشهر	۲/۳۵	۲/۵۶	۱/۹۹	۲/۳۰	۰/۹۱
۸	تهران	۱/۷۶	۱/۶۶	۱/۱۶	۱/۷۰	۰/۷۲
۹	چهارمحال و بختیاری	۲/۶	۲/۳۳	۲/۳۲	۲/۷۲	۱/۲۴
۱۰	خراسان جنوبی	۱/۸۱	۱/۱۴	۱/۰۹	۱/۷۶	۰/۷۸
۱۱	خراسان رضوی	۲/۷۶	۲/۸۲	۲/۴۶	۲/۸۰	۱/۴۲
۱۲	خراسان شمالی	۲/۰۲	۲/۳۲	۱/۹۱	۱/۹۶	۱/۸۴
۱۳	خوزستان	۲/۸۴	۳/۰۳	۲/۳۳	۲/۹۲	۱/۱۱
۱۴	زنجان	۲/۷۲	۲/۴۱	۱/۷۷	۲/۸۰	۰/۹۴
۱۵	سمنان	۲/۲۷	۲/۳۲	۲/۱۱	۲/۲۳	۱/۲۵
۱۶	سیستان و بلوچستان	۱/۴۲	۱/۷۵	۱/۳۳	۱/۵۰	۰/۳۲
۱۷	فارس	۲/۹۴	۳/۰۹	۲/۸۹	۲/۹۸	۱/۲۵
۱۸	قزوین	۳/۱۵	۳/۴۸	۲/۱۳	۳/۲۰	۰/۶۱
۱۹	قم	۲/۹۳	۲/۸۶	۱/۹۸	۲/۸۶	۱/۸۸
۲۰	کردستان	۲/۸۱	۲/۸۱	۱/۲۵	۲/۸۰	۰/۵۲
۲۱	کرمان	۱/۷۳	۱/۶۹	۱/۲۵	۱/۷۷	۱/۰۹
۲۲	کرمانشاه	۲/۶۳	۲/۳۲	۱/۷۹	۲/۶۵	۱/۶۷
۲۳	کهگیلویه و بویراحمد	۱/۰۴	۱/۳۷	۱/۰۵	۱/۰۵	۰/۹۷
۲۴	گلستان	۱/۲۴	۱/۷۵	۱/۶۶	۱/۵۱	۱/۵۳
۲۵	گیلان	۷/۳۸	۶/۷۵	۱/۸۶	۷/۵۴	۱/۵۶
۲۶	لرستان	۲/۴۷	۲/۸۰	۱/۴۱	۲/۶۷	۰/۸۹
۲۷	مازندران	۲/۹۷	۲/۷۳	۱/۳۰	۲/۸۱	۰/۶۳
۲۸	مرکزی	۲/۰۳	۲/۱۳	۱/۳۶	۲/۱۸	۱/۲۴
۲۹	هرمزگان	۲/۳۶	۲/۲۰	۱/۶۱	۲/۳۰	۰/۸۴
۳۰	همدان	۳/۹۸	۳/۶۷	۲/۵۹	۳/۹۳	۰/۹۵
۳۱	یزد	۲/۳۷	۲/۴۵	۱/۲۵	۲/۴۰	۱/۱۱

برای مدیریت مؤثر تولید و توزیع انرژی، فهم ماهیت احتمالاتی خروجی انرژی ضروری است. کمیت‌هایی از قبیل احتمال بیشتر شدن یا کمتر شدن توان خروجی انرژی از یک مقدار آستانه^۲ جذابیت ویژه‌ای برای بررسی خواهد داشت. با این انگیزه پیش‌گویی‌های سرعت باد را برای پیش‌گویی توان خروجی انرژی مورد استفاده قرار داده و در جدول ۳ ارائه شده است. توان خروجی انرژی باد به صورت $P = \frac{593}{\rho} \rho u^3$ برآورد می‌شود، که در آن P بیانگر توان خروجی انرژی باد بر

حسب $\rho = 1/227, W/m^2$ بر حسب kg/m^3 چگالی هوا و u سرعت آنی باد در یک ایستگاه خاص بر حسب متر بر ثانیه است. برای خروجی‌های انرژی، آستانه منتخب دلخواه $120 W/m^2$ و $240 W/m^2$ ، احتمال‌های خروجی انرژی واقعی کمتر از ۱۲۰ و بیشتر از ۲۴۰ محاسبه می‌شود که کمترین و بیشترین آستانه یک توربین برای توان خروجی است. برآوردهای هفتگی این احتمال‌ها در جدول ۳ ارائه شده است. برای افزایش ظرفیت تولید برق بادی در ایستگاه‌هایی که توان خروجی آن‌ها زیاد است می‌توان مدیریت صحیح و مؤثری برای آن‌ها به کار گرفت.

جدول ۳: برآورد احتمال وقوع متوسط خروجی انرژی هفتگی در پنج هفته اول سال ۲۰۱۶

ردیف	نام ایستگاه	$P_r(120 < Power < 240)$	ردیف	نام ایستگاه	$P_r(120 < Power < 240)$
۱	آذربایجان شرقی	۰/۳۳	۱۶	سیستان و بلوچستان	۰/۵۵
۲	آذربایجان غربی	۰/۶۱	۱۷	فارس	۰/۲۵
۳	اردبیل	۰/۷۲	۱۸	قم	۰/۵
۴	اصفهان	۰/۳	۱۹	کردستان	۰/۵۰
۵	البرز	۰/۸۷	۲۰	کرمان	۰/۶۶
۶	ایلام	۰/۷۱	۲۱	کرمانشاه	۰/۹۱
۷	بوشهر	۰/۶۸	۲۲	کهگیلویه و بویراحمد	۰/۴
۸	تهران	۰/۷	۲۳	گلستان	۰/۳۳
۹	چهارمحال و بختیاری	۰/۷۸	۲۴	گیلان	۰/۶
۱۰	خراسان جنوبی	۰/۵۱	۲۵	لرستان	۰/۷۰
۱۱	خراسان رضوی	۰/۴۸	۲۶	مازندران	۰/۴
۱۲	خراسان شمالی	۰/۲۱	۲۷	مرکزی	۰/۶۸
۱۳	خوزستان	۰/۷۷	۲۸	هرمزگان	۰/۷۴
۱۴	زنجان	۰/۷۱	۲۹	همدان	۰/۲۶
۱۵	سمنان	۰/۴۱	۳۰	یزد	۰/۰۶

بحث و نتیجه‌گیری

برای پیشگویی فضایی-زمانی در یک میدان تصادفی نامانا و افزایش سرعت محاسبات، از تانسور کوواریانس فضایی-زمانی مبتنی بر دو تجزیه تاکر و CP مورد بررسی قرار گرفت و نشان داده شد خطاهای محاسباتی کاهش، دقت پیشگویی‌ها افزایش و سرعت محاسبات افزایش می‌یابد. بعلاوه نشان داده شد پیشگویی‌های مبتنی بر تجزیه تاکر از دقت بیشتری نسبت به تجزیه CP برخوردار هستند.

مراجع

- محمدزاده، م.، (۱۳۹۸)، آمار فضایی و کاربردهای آن، چاپ سوم، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران،
 عصمتی، م و محمدزاده، م. (۱۳۹۶)، پیش‌گویی فضایی-زمانی میدان‌های تصادفی نامانا و تفکیک‌ناپذیر، دومین سمینار
 آمار فضایی و کاربردهای آن، دانشگاه صنعتی شاهرود.

Gneiting, T. (2002), Nonseparable, Stationary Covariance Functions for Space-Time Data, *Journal of the American Statistical Association*, **97**, 590–600.

- Hitchcock, F. L. (1927). The Expression of a Tensor or a Polyadic as a Sum of Products, *Journal of Mathematical Physics*, **6**, 164–189.
- Kolda T.G. and Bader B.W. (2009), Tensor Decompositions and Applications, *SIAM Review*, **51**, 455–500.
- Kruskal, J. B. (1977), Three-Way Arrays: Rank and Uniqueness of Trilinear Decompositions, with Application to Arithmetic Complexity and Statistics, *Linear Algebra and its Applications*, **18**, 95–138.
- Lynch, P. (2008), The Origins of Computer Weather Prediction and Climate Modeling, *Journal of Computational Physics*, **227**, 3431–3444.
- Salcedo-Sanz S., Perez-Bellido, A.M., Ortiz-Garca, E.G., Portilla-Figueras, A., Prieto, L. and Correo, F. (2009), Accurate Short-Term Wind Speed Prediction by Exploiting Diversity in Input Data Using Banks of Artificial Neural Networks, *Neurocomputing*, **72**, 1336–1341.
- Soman, S., Zareipour, H., Malik, O. and Mandal, P. (2010), A Review of Wind Power and Wind Speed Forecasting Methods with Different Time Horizons, *North American Power Symposium (NAPS)*, Arlington, TX, 26–28 September, 1–8.
- Tucker, L. R. (1963) Implications of Factor Analysis of Three-Way Matrices for Measurement of Change, in Problems in Measuring Change, C. W. Harris, ed., *University of Wisconsin Press*, 122–137.



مقایسه کارائی آماری و محاسباتی روش‌های برآوردیابی در مدل‌های نیمه‌طیفی فضایی-زمانی

علی شهنواز^۱، علی محمدیان مصمم^۲، محمد حسن بهزادی^۳

^۱گروه آمار، دانشگاه آزاد اسلامی، واحد زنجان

^۲گروه آمار، دانشگاه زنجان

^۳گروه آمار، دانشگاه آزاد اسلامی، واحد علوم و تحقیقات، تهران

چکیده: در این مقاله، ابتدا مبانی نظری مدل‌سازی نیمه‌طیفی بررسی و به توصیف چند خاصیت از مدل‌های نیمه‌طیفی اخیر پرداخته می‌شود. سپس یک روش برای برآورد تابع کوواریانس فضایی-زمانی در حالت نیمه‌طیفی پیشنهاد شده است. به منظور ارزیابی عملکرد مدل‌های نیمه‌طیفی ارائه شده، دو شبیه‌سازی انجام و در هرکدام از آن‌ها روش برآورد پیشنهادی با سایر روش‌ها مقایسه شده است. روش مورد نظر عملکرد بهتری نسبت به سایر روش‌ها به خصوص برای مجموعه مشاهدات بزرگ دارد. سرانجام برای تحلیل داده‌های واقعی مربوط به متوسط روزانه سرعت باد در استان زنجان از روش پیشنهادی استفاده شده است.

واژه‌های کلیدی: کوواریانس فضایی-زمانی، مدل فضایی-زمانی، مدل نیمه‌طیفی، درستمایی وایتل.
کد موضوع بندی ریاضی (۲۰۱۰): 62M30، 62M15.

۱ مقدمه

در آمار کلاسیک معمولاً داده‌ها ناهمبسته فرض می‌شوند ولی در عمل موارد زیادی وجود دارند که این فرض ناهمبستگی نقض شده است. به عنوان مثال در سری‌های زمانی **شاموای و استافر (۲۰۱۷)** مشاهدات دارای نوعی همبستگی از طریق زمان رخداد آنها هستند. همچنین در علوم مربوط به زمین آمار (**چیلز و دلفنر، ۲۰۱۲**)، محیط زیست، هواشناسی، همه‌گیرشناسی و اقیانوس‌شناسی ملاحظه می‌شود که داده‌ها دارای نوعی همبستگی از نظر موقعیت مشاهدات هستند. به چنین

¹Spatial Data

^۱نام و ایمیل ارائه دهنده مقاله: علی شهنواز، shahnavaz.ali2000@yahoo.com

داده‌هایی که همبستگی آنها ناشی از موقعیت جغرافیایی است، داده‌های فضایی^۱ گفته می‌شود (کرسی و ویکل، ۲۰۱۳). داده‌هایی که هم از نظر موقعیت فضایی و هم از نظر زمانی وابسته باشند، داده‌های فضایی-زمانی^۲ نامیده می‌شوند. برای تحلیل سری‌های زمانی و داده‌های فضایی معمولاً دو استراتژی وجود دارد: یکی تحلیل سری‌های زمانی و داده‌های فضایی در قلمرو مشاهدات و دیگری تحلیل سری‌های زمانی و داده‌های فضایی در قلمرو فرکانس. در دو دهه اخیر مطالعات زیادی برای تعیین ساختار همبستگی، مدل‌بندی و تحلیل داده‌ها انجام شده است. تحلیل این گونه داده‌ها به روش اول مستلزم تعیین ساختار همبستگی به ترتیب، زمانی، فضایی و فضایی-زمانی از طریق تابع کواریانس آنها است. این تابع نقش به سزایی در پیش‌گویی موقعیت‌های فضایی یا زمانی فاقد مشاهده دارد. تحلیل داده‌ها به روش دوم مستلزم تعیین تابع چگالی طیفی زمانی، فضایی و فضایی-زمانی است. توابع کواریانس و توابع چگالی طیفی زمانی و فضایی محض بسیار زیادی در متون آماری وجود دارند (چپلس و دلفنر، ۲۰۱۲؛ کرسی، ۱۹۹۳). ولی برای ساخت توابع کواریانس فضایی-زمانی مطالعات علمی کمتری انجام گرفته است. یکی از ویژگی‌های مهم تابع کواریانس، ویژگی معین مثبت^۳ است و تابعی که چنین شرطی را داشته باشد، تابع کواریانس معتبر نامیده می‌شود.

ساخت توابع کواریانس فضایی-زمانی معتبر تا حدودی امری مشکل و در طی چند سال گذشته یکی از موضوعات علمی آمارشناسان دنیا بوده است. به این منظور برخی فرض‌های ساده‌کننده نظیر، مانایی^۴، همسانگردی^۵، کاملاً متقارن^۶ و تفکیک‌پذیری^۷ مورد استفاده قرار می‌گیرد (کرسی، ۱۹۹۳). یکی از روش‌ها برای ساخت توابع کواریانس فضایی-زمانی معتبر استفاده از نمایش طیفی^۸ تابع کواریانس است که بر اساس قضیه بوخنر (۱۹۵۵) تعیین می‌شود. این قضیه به درستی تحلیل در قلمرو مشاهدات را به تحلیل در قلمرو فرکانس مرتبط می‌سازد. بوخنر نشان داد که یک تابع پیوسته، معین مثبت است اگر و تنها اگر تبدیل فوریه‌ی یک اندازه نامنفی متناهی باشد. از آن‌جا که توابع کواریانس فضایی-زمانی هر میدان تصادفی مانای مرتبه دوم، موجود و معین مثبت است، براساس قضیه بوخنر رده توابع کواریانس فضایی-زمانی مانا روی $R^d \times R$ معادل رده تبدیل فوریه‌ی اندازه‌های نامنفی و متناهی روی $R^d \times R$ است، لذا با توجه به این قضیه، نمایش طیفی تابع کواریانس بر اساس تابع چگالی طیفی عبارت است از:

$$K(s, t) = \iint g(\lambda, \omega) \exp\{is'\lambda + it\omega\} d\lambda d\omega, \quad (s, t) \in R^d \times R \quad (1.1)$$

که در آن s و t به ترتیب تاخیرهای فضایی و زمانی و همچنین $\lambda \in R^d$ و $\omega \in R$ فرکانس‌های فضایی و زمانی هستند و نماد " " نشانگر ترانهاده‌ی یک بردار است. بنابراین، با اختیار کردن تابع چگالی طیفی می‌توان یک تابع کواریانس فضایی-زمانی معتبر بنا نهاد.

در تحلیل داده‌های فضایی-زمانی علاوه بر تحلیل در قلمرو مشاهدات و قلمرو فرکانس می‌توان قلمرو جدیدی در مابین این دو تعریف کرد که به اصطلاح نیمه‌طیفی^۹ نامیده می‌شود. از مهمترین مزایای این روش استفاده از انتگرال صرفاً فضایی به جای انتگرال چندگانه فضایی-زمانی در رابطه (۱.۱) است. کرسی و هوانگ (۱۹۹۹) با فرض انتگرال‌پذیری تابع کواریانس و به کارگیری قضیه بوخنر، روشی برای ساخت کواریانس‌های تفکیک‌ناپذیر مانا ارائه کردند که روش آنان محدود به رده‌ای از توابع می‌شود که انتگرال d -گانه تبدیل فوریه آنها به صورت تحلیلی قابل حل باشد (کاملر، ۲۰۰۷). گنتینگ (۲۰۰۲) برای غلبه بر انتگرال صرفاً فضایی فوق با استفاده از توابع کاملاً یکنوا و برنشتاین رده‌ای از مدل‌های

²Spatio-Temporal Data

³Positive Definite

⁴Stationary

⁵Isotropic

⁶Completely Symmetric

⁷Separable

⁸Spectral Representation

⁹Half Spectral

کواریانس فضایی-زمانی تفکیک‌ناپذیر ارائه کرد. علیرغم سادگی بسیار زیاد روش گنیتینگ، تابع کواریانس حاصل یک تابع کواریانس کاملاً متقارن است. علاوه بر آن **کنت و همکاران (۲۰۱۱)** نشان دادند که تابع کواریانس ارائه شده توسط گنیتینگ دارای مشکل گودی^{۱۰} هست، به این معنی که به ازای برخی مقادیر تاخیر فضایی، تابع کواریانس زمانی مربوطه به صورت یک‌نوا نزولی نیست. **امیدی و محمدزاده (۲۰۱۵)**، **امیدی و محمدزاده (۱۳۹۲)** با استفاده از توابع مفصل برای حل مشکل گودی راهکاری ارائه نموده‌اند. **هورل و استاین (۲۰۱۷)** با استفاده از نمایش نیمه‌طیفی تابع کواریانس فضایی-زمانی مدل‌های گاوسی فضایی-زمانی جدیدی را ارائه کرده‌اند.

در این مقاله چند مدل نیمه‌طیفی معرفی و از بین آنها بهترین مدل برای برازش به داده‌های میانگین سرعت باد روزانه استان زنجان انتخاب و با استفاده از روش ماکسیمم درست‌نمایی چند متغیره وایتل، پارامترهای مدل انتخابی برآورد شده‌اند.

۲ مدل‌های نیمه‌طیفی

نمایش نیمه‌طیفی تابع کواریانس $K(s, t)$ را می‌توان با انتخاب ترتیب انتگرال‌گیری نسبت به λ یا ω به صورت نیمه‌طیفی زمانی یا نیمه‌طیفی فضایی بدست آورد. در این مقاله نمایش نیمه‌طیفی در زمان به صورت زیر مورد بحث قرار خواهد گرفت.

$$K(s, t) = \int_R H(s, \omega) \exp(it\omega) d\omega,$$

که در آن:

$$H(s, t) = f(\omega) \mathbb{C}(s\delta(\omega)) \exp(i\theta(\omega)\phi'(s)), \quad (1.2)$$

که در آن f تابع چگالی طیفی صرفاً زمانی، $\mathbb{C}$ تابع همبستگی معتبر در R^d با یک تابع چگالی طیفی، δ تابعی زوج و مثبت که نشان‌دهنده اثر تعامل فضایی-زمانی، $\theta(\cdot)$ تابعی حقیقی و فرد و ϕ برداریکه در R^d است (**استاین، ۲۰۰۵**).

۱.۲ چند خاصیت از مدل‌های نیمه‌طیفی

مدل معرفی شده در رابطه (۱.۲) در حالت کلی نامتقارن است. در این مقاله ابتدا حالت فضایی-زمانی کاملاً متقارن این مدل یعنی حالتی که $\theta(\omega) = 0$ در نظر گرفته می‌شود. با این فرض، مدل نیمه‌طیفی رابطه (۱.۲) به صورت $f(\omega) \mathbb{C}(s\delta(\omega))$ است. **هورل و استاین (۲۰۱۷)** نشان دادند که بسیاری از خواص مدل‌سازی $K(s, t)$ مستقل از θ یا ϕ حفظ می‌شود. خاصیت تعامل فضایی-زمانی با استفاده از δ تعیین می‌شود. لازم به ذکر است وقتی δ مقدار ثابت باشد توابع کواریانس حاشیه‌ای فضایی و زمانی مستقل از هم تعیین می‌شوند. در این حالت $K(s, t)$ متناسب با حاصل ضرب توابع کواریانس فضایی و زمانی است یعنی

$$K(s, t) \propto K(s, 0)K(0, t)$$

این مدل‌ها را مدل‌های تفکیک‌پذیر (**رودریگاز و مجیا، ۱۹۷۴**) گویند که به علت عدم تطابق با واقعیت داده‌ها، مورد انتقاد هستند. همان‌طور که در بخش بعد مشاهده خواهد شد در مدل‌های نیمه‌طیفی در تابع کواریانس گودی ظاهر می‌شود (**کنت و همکاران، ۲۰۱۱**). به این معنی که اگر فرض کنیم s_{here} نشان‌دهنده موقعیت جاری و t_{now} نیز زمان جاری باشد و همبستگی $Z(s_{here}, t_{now})$ را با مقدار فرایند در یک همسایگی از موقعیت s_{there} در دو زمان ممکن $t_{yesterday}$ و t_{now}

¹⁰Dimple

در نظر بگیریم، در این حالت همبستگی با $Z(s_{there}, t_{yesterday})$ یا معادل آن با $Z(s_{there}, t_{tomorrow})$ ممکن است بزرگ‌تر یا کوچک‌تر از همبستگی با $Z(s_{there}, t_{now})$ باشد.

استاین (۲۰۱۱) شرط زیر را برای تابع چگالی طیفی فضایی-زمانی g ، به عنوان یکی از شرایط مطلوب برای توابع کاملاً طیفی، ذکر کرده است. با برقراری این شرط، خاصیت غربالگری Effect Screening در این مدل‌ها تضمین می‌شود (ژورنل و هوجبرگتس، ۱۹۷۸).

$$\lim_{(\lambda, \omega) \rightarrow \infty} \sup_{(u, v) < R} \left| \frac{g(\lambda + \nu, \omega + u)}{g(\lambda, \omega)} - 1 \right| = 0, \quad (2.2)$$

که در آن R می‌تواند هر مقدار متناهی باشد. لازم به ذکر است که مدل نیمه‌طیفی رابطه (۱.۲) در حالت کلی در شرط (۲.۲) صدق نمی‌کند. فرض کنید h چگالی طیفی تابع کواریانس فضایی $\mathbb{C}$ باشد پس نمایش کاملاً طیفی مدل استاین (۲۰۰۵) به صورت زیر است (مصمم و کنت، ۲۰۱۶):

$$K(s, t) = \iint_{R^d R} f(\omega) \delta(\omega)^{-d} h\left(\frac{\lambda - \theta(\omega)\phi}{\delta(\omega)}\right) \exp(is'\lambda + it\omega) d\lambda d\omega.$$

در حالتی که $g(\lambda, \omega)$ یک تابع چگالی طیفی کاملاً متقارن در نظر گرفته شود خواهیم داشت:

$$g(\lambda, \omega) = f(\omega) \delta(\omega)^{-d} h\left(\frac{\lambda}{\delta(\omega)}\right). \quad (3.2)$$

بنابراین نمایش کلی g را می‌توان به صورت $g(\lambda - \theta(\omega)\phi, \omega)$ نوشت. در قضیه ۱.۲ شرایط لازم برای اینکه $g(\lambda, \omega)$ یا $g(\lambda - \theta(\omega)\phi, \omega)$ در شرط (۲.۲) صدق کند بیان شده است.

قضیه ۱.۲ (هول و استاین، ۲۰۱۷). فرض کنید $g(\lambda, \omega)$ تابع چگالی طیفی تابع کواریانس فضایی-زمانی $K(s, t)$ باشد که در آن λ فرکانس فضایی و ω فرکانس زمانی است.

الف. فرض کنید $\theta(\cdot)$ به صورت موضعی کراندار باشد. چگالی طیفی $g(\lambda, \omega)$ در (۲.۲) صدق می‌کند اگر و فقط اگر $g(\lambda - \theta(\omega)\phi, \omega)$ در (۲.۲) صدق کند.

ب. فرض کنید $g(\lambda, \omega)$ اکیداً مثبت باشد و در (۲.۲) صدق کند. اگر $g(\lambda - \theta(\omega)\phi, \omega)$ در (۲.۲) صدق کند، آنگاه باید $\theta(\omega)$ به صورت موضعی کراندار باشد.

در واقع چون استراتژی‌های ساخت مدل اغلب با مدل‌های کاملاً متقارن شروع می‌شوند، به کمک این قضیه، استراتژی ساخت مدل را می‌توان به مدل‌های نامتقارن تعمیم داد. حال با این ذهنیت از قضیه ۱.۲ بر روی مدل‌های نیمه‌طیفی کاملاً متقارن تمرکز خواهد شد. هول و استاین هول و استاین (۲۰۱۷) نشان دادند برای اینکه $g(\lambda, \omega)$ در شرط (۲.۲) صدق کند باید طیف h و تابع δ به صورت

$$h\left(\frac{\lambda}{\delta(\omega)}\right) = H(\lambda, \omega) \delta(\omega)^d / f(\omega),$$

قابل تجزیه باشند که در آن H در شرط (۲.۲) صدق می‌کند. بنابراین نمایش ساده‌تری به صورت $g(\lambda, \omega) = p(\omega)q(\lambda, \omega)$ می‌توان برای g در نظر گرفت. در ادامه شرایطی را که p و q باید داشته باشند تا g در شرط (۲.۲) صدق کند ذکر خواهد شد.

محدودیت ۱.۲. فرض کنید نمایش طیفی تابع کواریانس K شکلی به صورت $g(\lambda, \omega) = p(\omega)q(\lambda, \omega)$ داشته باشد در اینصورت دو گزاره زیر برای چگالی‌های طیفی که به صورت فوق پارامتری شده‌اند برقرار هستند:

الف. اگر g و q در شرط (۲.۲) صدق کنند آنگاه باید p نسبت به q ثابت باشد.

ب. اگر g در شرط (۲.۲) صدق کند، آنگاه برای هر نقطه ثابت ω که $p(\omega) > 0$ متناهی و مثبت باشد، طیف حاشیه‌ای $g_{\omega^*}(\lambda) = g(\lambda, \omega)$ باید در شرط (۲.۲) صدق کند.

پیامد این محدودیت این است که برای تعیین مدل‌هایی که در شرط (۲.۲) صدق می‌کنند، تابع δ را نمی‌توان مستقل از f و $\mathbb{C}$ در نظر گرفت.

محدودیت ۲.۲. فرض کنید $g(\lambda, \omega) = p(\omega)q(\lambda/\delta(\omega))$ که در آن $p(\omega)$ برای تمام $|\omega|$ ‌های به اندازه کافی بزرگ مثبت و متناهی است، q یک تابع انتگرال‌پذیر، پیوسته و نامنفی و $\delta(\omega)$ زوج با حد خوش تعریف وقتی که $|\omega| \rightarrow \infty$ اگر در (۲.۲) صدق کند آنگاه $\lim_{|\omega| \rightarrow \infty} \delta(\omega) = \infty$.

محدودیت‌های ۱.۲ و ۲.۲ شکل δ و h را تعیین می‌کنند.

۲.۲ رده‌ای از مدل‌های نیمه‌طیفی

در این بخش در قالب چند مثال، نمونه‌هایی از مدل‌های نیمه‌طیفی کاملاً متقارن معرفی می‌شوند.

مثال ۱.۲. مدل تفکیک‌پذیر (کرسی و هوانگ، ۱۹۹۹) با تابع کواریانس حاشیه‌ای نمایی: فرض کنید $f(\omega) = (a^2 + \omega^2)^{-1} \exp\{-b \|s\|\}$ تابع چگالی طیفی زمانی باشد که تابع کواریانس فضایی آن به صورت $\mathbb{C}(s\delta(\omega)) = \exp\{-b \|s\|\}$ است. همچنین بردار پارامترها به صورت $\theta = (a, b, \sigma^2)'$ را در نظر بگیرید که در آن $a \geq 0$ پارامتر مقیاس زمان σ^2 ، $b \geq 0$ پارامتر مقیاس فضا σ^2 و σ^2 واریانس میدان تصادفی (ازاره) است. در اینصورت مدل نیمه‌طیفی متناظر با آن به صورت

$$H(s, \omega) = f(\omega)\mathbb{C}(s\delta(\omega)) = \sigma^2 (a^2 + \omega^2)^{-1} \exp\{-b \|s\|\} \quad (4.2)$$

خواهد بود و تابع کواریانس فضایی-زمانی مربوطه به صورت

$$K(s, t) = \sigma^2 \exp(-a|t| - b \|s\|)$$

است. با استفاده از رابطه (۳.۲) تابع چگالی طیفی فضایی-زمانی به صورت زیر خواهد بود:

$$g(\lambda, \omega) = \sigma^2 \left(\frac{2a}{a^2 + \omega^2} \right) \left(\frac{2b}{b^2 + \|\lambda\|^2} \right),$$

لازم به ذکر است که در این حالت $\delta(\omega)$ ثابت، تابع کواریانس و همچنین تابع چگالی طیفی تفکیک‌پذیر و $H(s, \omega)$ تابعی حقیقی مقدار است.

مثال ۲.۲. مدل تفکیک‌پذیر (کرسی و هوانگ، ۱۹۹۹) با توابع کواریانس حاشیه‌ای مربع نمایی فضایی و نمایی زمانی

¹¹Time Scale Parameter

¹²Space Scale Parameter

: فرض کنید $f(\omega) = 2a\sigma^2(a^2 + \omega^2)^{-1}$ تابع چگالی طیفی زمانی باشد که تابع کواریانس فضایی آن به صورت $\delta(\omega) = 1, \mathbb{C}(s\delta(\omega)) = \exp\{-b \|s\|^2\}$ است. همچنین بردار پارامترها به صورت $\theta = (a, b, \sigma^2)'$ را در نظر بگیرید که در آن $a \geq 0$ پارامتر مقیاس زمان، $b \geq 0$ پارامتر مقیاس فضا و σ^2 واریانس میدان تصادفی (ازاره) است. در اینصورت مدل نیمه‌طیفی متناظر با آن به فرم

$$H(s, \omega) = f(\omega)\mathbb{C}(s\delta(\omega)) = 2a\sigma^2(a^2 + \omega^2)^{-1} \exp\{-b \|s\|^2\}$$

خواهد بود و تابع کواریانس فضایی-زمانی مربوطه به صورت

$$K(s, t) = \sigma^2 \exp(-a|t| - b \|s\|^2)$$

است. با استفاده از رابطه (۳.۲) تابع چگالی طیفی فضایی-زمانی به صورت زیر خواهد بود:

$$g(\lambda, \omega) = \frac{2a\sigma^2\sqrt{\pi}}{b(a^2 + \omega^2)} \exp\left(-\frac{\|\lambda\|^2}{4b^2}\right).$$

لازم به ذکر است که در این حالت $\delta(\omega)$ ثابت، تابع کواریانس و همچنین تابع چگالی طیفی تفکیک‌پذیر و $H(s, \omega)$ تابعی حقیقی مقدار است.

مثال ۳.۲. مدل تفکیک‌ناپذیر کاملاً متقارن (گنتینگ، ۲۰۰۲): فرض کنید

$$H(s, \omega) = \frac{\sigma^2}{(b\|s\|)^\alpha + 1} \times \frac{2a((b\|s\|)^\alpha + 1)^{-\beta/2}}{a^2((b\|s\|)^\alpha + 1)^{-\beta} + \omega^2}, \quad (5.2)$$

که در آن بردار پارامترها $\theta = (a, b, \alpha, \beta, \sigma^2)'$ ، $a \geq 0$ پارامتر مقیاس زمان، $b \geq 0$ پارامتر مقیاس فضا، σ^2 واریانس میدان تصادفی (ازاره)، α پارامتر همواری Smoothness Parameter در بازه $[0, 2]$ و β پارامتر تعامل^{۱۳} فضایی-زمانی است، در اینصورت تابع کواریانس فضایی-زمانی متناظر با آن به صورت زیر خواهد بود:

$$K(s, \omega) = \frac{\sigma^2}{(b\|s\|)^\alpha + 1} \times \exp\left(\frac{-a|t|}{((b\|s\|)^\alpha + 1)^{\beta/2}}\right).$$

متأسفانه در این حالت نمایش بسته‌ای برای $g(\lambda, \omega)$ وجود ندارد و همچنین نمی‌توان آن را به صورت مدل استاین یعنی $f(\omega)\mathbb{C}(s\delta(\omega))$ نشان داد و لذا محاسبات به صورت عددی انجام می‌شود.

مثال ۴.۲. مدل تفکیک‌ناپذیر کاملاً متقارن (هول و استاین، ۲۰۱۷): فرض کنید $h(\lambda) = \phi(\alpha^2 + \|\lambda\|^2)^{-(v+d/2+1/2)}$ چگالی طیفی ماترن با پارامتر همواری $v + \frac{1}{p} > 0$ ، پارامتر دامنه وارون $\alpha > 0$ و پارامتر مقیاس $\phi > 0$ باشد. برای یک f مفروض به منظور اطمینان از عدم نقض محدودیت ذکر شده، $\delta(\omega)$ به صورت زیر در نظر گرفته شده است:

$$\delta(\omega) = f(\omega)^{-\frac{1}{v+1/2}}, \quad (6.2)$$

با جایگذاری رابطه (۶.۲) در (۳.۲) خواهیم داشت:

$$g(\lambda, \omega) = \phi \left(\alpha^2 f(\omega)^{-\frac{1}{v+1/2}} + \|\lambda\|^2 \right)^{-(v+d/2+1/2)}. \quad (7.2)$$

¹³Separability Parameter

اگر f به صورت $f \sim c_f \omega^{-k}$ در نظر گرفته شود برای بعضی از مقادیر $k > 1$ و c_f که یک مقدار ثابت مثبت وابسته به f می‌باشد، آنگاه $g(\lambda, \omega)$ تعریف شده به صورت فوق، رده‌ای از مدل‌ها را توصیف می‌کند که در شرط (۲.۲) صدق می‌کنند. نمایش نیمه‌طیفی رابطه (۷.۲) با جایگذاری تابع کواریانس ماترن در رابطه (۱.۲) به صورت

$$\begin{aligned} H(\mathbf{s}, \omega) &= f(\omega) \mathbb{C}(\mathbf{s} \delta(\omega)) \\ &= f(\omega) \frac{\phi \pi^{d/2}}{2^{v-1/2} \Gamma(v + (d+1)/2) \alpha^{2v+1}} \left(\alpha \|\mathbf{s}\| f(\omega)^{-1/2 \frac{1}{v+1/2}} \right)^{v+1/2} \\ &\quad \times K_{v+1/2} \left(\alpha \|\mathbf{s}\| f(\omega)^{-1/2 \frac{1}{v+1/2}} \right), \end{aligned} \quad (۸.۲)$$

حاصل می‌شود که در آن $K_{v+1/2}$ تابع بسل مرتبه دوم است. اگر مقادیری که به ω بستگی ندارند در ϕ خلاصه کنیم، با فرض

$$\mathcal{M}_{V+1/2}(\mathbf{s}) = \|\mathbf{s}\|^{v+1/2} K_{v+1/2}(\|\mathbf{s}\|),$$

رابطه (۸.۲) را می‌توان به صورت زیر نوشت:

$$H(\mathbf{s}, \omega) = f(\omega) \mathbb{C}(\mathbf{s} \delta(\omega)) = \phi f(\omega) \mathcal{M}_{V+1/2} \left(\alpha \|\mathbf{s}\| f(\omega)^{-1/2 \frac{1}{v+1/2}} \right). \quad (۹.۲)$$

مثال ۵.۲. ماترن زمانی (هول و استاین، ۲۰۱۷): فرض کنید $f(\omega) = (\beta^2 + \omega^2)^{k+1/2}$ چگالی طیفی ماترن یک بعدی باشد که برای سادگی، پارامتر مقیاس برابر یک فرض شده است. با جایگذاری آن در (۶.۲) و (۷.۲) نمایش طیفی و نیمه‌طیفی به صورت زیر خواهد بود:

$$g(\lambda, \omega) = \phi \left(\alpha^2 (\beta^2 + \omega^2)^{(k+1/2)/(v+1/2)} + \|\lambda\|^2 \right)^{-(v+d/2+1/2)}.$$

$$f(\omega) \mathbb{C}(\mathbf{s} \delta(\omega)) = \phi \left((\beta^2 + \omega^2)^{-(k+1/2)} \right) \mathcal{M}_{v+1/2} \left(\alpha \|\mathbf{s}\| (\beta^2 + \omega^2)^{1/2 \frac{k+1/2}{v+1/2}} \right). \quad (۱۰.۲)$$

لازم به ذکر است به منظور بهبود مدل‌های استفاده شده همانند گنتینگ (۲۰۰۲) و استاین (۲۰۰۵) به هر پنج مدل مورد مطالعه یک اثر قطعه‌ای^{۱۴} اضافه شده است. بنابراین نمایش نیمه‌طیفی مدل‌های مورد بحث به صورت

$$f(\omega) [\mathbb{C}(\mathbf{s} \delta(\omega)) + \eta^2 \mathbb{I}_{\mathbf{s}=\mathbf{0}}],$$

خواهد بود که در آن $\mathbb{I}_{\mathbf{s}=\mathbf{0}}$ تابع نشانگر و $\eta^2 \geq 0$ اثر قطعه‌ای است.

۳ روش‌های برآورد پارامترها

در این بخش به معرفی چند روش برای برازش مدل‌های معرفی شده در بخش قبل پرداخته و سپس مدل‌های پیشنهادی در بخش قبل به داده‌های میانگین روزانه سرعت باد استان زنجان برازش شده است. چهار روش مختلف ماکسیم درستمایی (ML)، درستمایی ترکیبی (CL)، حداقل مربعات وزنی (WLS) و درستمایی وایتل^{۱۵} (WL) برای مدل‌های معرفی شده بکار رفته است.

^{۱۴}Nugget Nugget

^{۱۵}Whittle Likelihood

برای برآورد پارامترها وقتی که توزیع مشاهدات معلوم است از روش درست‌نمایی ماکسیمم، استفاده می‌شود. به عنوان مثال اگر میدان تصادفی از توزیع گاوسی (ماردیا و مارشال، ۱۹۸۴) پیروی کند می‌توان از روش درست‌نمایی ماکزیمم برای برآورد تغییرنگار و هم تغییر نگارها استفاده کرد.

فرض کنید $\mathbf{Z} = (Z(s_1, t_1), \dots, Z(s_n, t_n))'$ ، n مشاهده از میدان تصادفی گاوسی با میانگین صفر و ماتریس کواریانس Σ باشد. در این صورت لگاریتم تابع درست‌نمایی به صورت

$$\ell(\theta) \propto \log |\Sigma(\theta)| - \mathbf{Z}'\Sigma(\theta)^{-1}\mathbf{Z},$$

است که در آن $\Sigma(\theta) = \text{Var}(\mathbf{Z})$ به $\theta \in \Theta \subseteq \mathbb{R}^p$ وابسته است. برای یک میدان تصادفی گاوسی با یک تابع کواریانس پارامتری، برای محاسبه دقیق درست‌نمایی به محاسبه معکوس و دترمینان ماتریس کواریانس نیاز است. هنگامی که تعداد مشاهدات زیاد است این ارزیابی مشکل خواهد بود. این واقعیت انگیزه‌ی اصلی تحقیق برای دست یافتن به تقریب درست‌نمایی است که نیاز به محاسبات کمتری داشته باشد.

در مواردی که توزیع مشاهدات معلوم نیست، برای برآورد پارامترهای مدل می‌توان از روش کم‌ترین توان‌های دوم وزنی استفاده کرد (کرسی، ۱۹۸۵). در این روش فقط از عناصر روی قطر ماتریس کواریانس Σ استفاده می‌شود و بیشترین وزن به فاصله‌های نزدیک و وزن‌های کم به فاصله‌های دور اختصاص داده می‌شوند. در این روش θ به نحوی انتخاب می‌شود که رابطه زیر مینیمم شود (کرسی، ۱۹۹۳):

$$WLS(\theta) = \sum_{l=1}^L \sum_{u=1}^U |N(h(l); u)| \left\{ \frac{\hat{\gamma}(h(l); u)}{\hat{\gamma}(h(l); u | \theta)} - 1 \right\}^2,$$

که در آن، $N(h(l); u) \equiv \{(i, j, t, t') : s_i - s_j \in \text{Tol}(h(l)); |t - t'| = u\}$ که در فاصله (h, u) از یکدیگر قرار دارند. در اینجا $\text{Tol}(h(l))$ ناحیه تحمل^{۱۶} اطراف $h(l)$ است. روش درست‌نمایی ترکیبی به طور کلی یک روش شبه درست‌نمایی است که از توزیع حاشیه‌ای یا شرطی با بعد کمتری ساخته شده است. ایده اصلی درست‌نمایی ترکیبی در ابتدا توسط بسیج (۱۹۷۴) برای مدل‌سازی داده‌های فضایی ارائه گردید و بعداً توسط لیندسی (۱۹۸۸) به عنوان درست‌نمایی ترکیبی معرفی شد. در مواردی که برآورد و استنباط برای پارامترها و مدل‌های شبه پارامتری که ارزیابی توابع درست‌نمایی آنها با توجه به ساختار پیچیده داده‌ها دشوار است، محبوبیت این روش بیشتر است. دو روش برای ساخت تابع درست‌نمایی ترکیبی مطرح است: روش درست‌نمایی حاشیه‌ای ترکیبی و روش درست‌نمایی شرطی ترکیبی. طبق تعریف لیندسی تابع درست‌نمایی ترکیبی به صورت

$$CL(\theta, z) = \prod_{i=1}^K L_k(\theta, z),$$

است که در آن $L_k = (\theta, z)$ ، $k = 1, \dots, K$ توابع درست‌نمایی حاشیه‌ای یا شرطی هستند.

درست‌نمایی وایتل (۱۹۵۴) یک تقریب برای درست‌نمایی سری‌های زمانی گاوسی مانا است. در یک مدل سری‌زمانی گاوسی مانا، تابع درست‌نمایی (بطور معمول گاوسی) توسط پارامترهای میانگین و کواریانس مربوط به مدل تعیین می‌شود و با وجود تعداد زیاد مشاهدات ممکن است ماتریس کواریانس بسیار بزرگ و محاسبات در عمل پرهزینه شود. بنابراین با توجه به مانایی، ماتریس کواریانس دارای یک ساختار نسبتاً ساده می‌شود که با استفاده از یک تقریب، محاسبات به طور قابل توجهی ساده خواهد شد. تابع درست‌نمایی وایتل چند متغیره به صورت

$$L(\theta) = \sum_{k=1}^{T-1} \left\{ \log |F(\omega_k, \theta)| + \text{tr} \left(F(\omega_k, \theta)^{-1} I(\omega_k) \right) \right\} \quad (1.3)$$

¹⁶Tolerance Region

است که در آن F و I به ترتیب ماتریس‌های طیفی و دوره نگار تجربی هستند.

۴ نتایج عددی

در این بخش یک مطالعه شبیه‌سازی با هدف ارزیابی عملکرد روش نیمه‌طیفی با توجه به مدل‌های کواریانس فضایی-زمانی تفکیک‌پذیر و تفکیک‌ناپذیر انجام شده است. در هرکدام از آنها روش پیشنهادی با چهار روش برآورد CL ، ML ، WLS و MWL و با تعداد مختلفی مشاهده مورد مقایسه قرار خواهد گرفت و سرانجام روش انتخاب شده برای داده‌های میانگین سرعت باد روزانه استان زنجان به کاررفته است. تعداد مشاهدات به گونه‌ای انتخاب شده که محاسبات مربوط به ماکسیمم درست‌نمایی به سادگی انجام شود. روش‌ها بر حسب میانگین مجذور خطاها (MSE) با ۱۰۰۰ تکرار از میدان تصادفی گاوسی $Z(s; t)$ با میانگین صفر و واریانس یک که در آن $t \in 1, \dots, T$ ، $T = 60, 90$ و $s_1 \in [0, 10]$ و $s_2 \in [0, 10]$ با $n = 10, 20$ و همچنین سرعت انجام محاسبات در یک تکرار مقایسه شده‌اند.

۱.۴ شبیه‌سازی با مدل تفکیک‌پذیر

ابتدا مدل تفکیک‌پذیر با تابع کواریانس حاشیه‌ای نمایی مثال ۱.۲ در نظر گرفته شده و نتایج برازش بر اساس ۱۰۰۰ تکرار در جدول ۱ نشان داده شده است. همانطوری که از جدول ۱ ملاحظه می‌شود MSE روش ML طبق انتظار کمترین مقدار و روش برآورد MWL نسبت به CL و WLS از نظر MSE بهتر عمل می‌کند. لازم به ذکر است که با افزایش n و یا MSE ، T برآوردها به کندی کاهش می‌یابد.

برای مدل تفکیک‌پذیر با تابع کواریانس حاشیه‌ای نمایی مثال ۱.۲ زمان لازم برای یک تکرار روش‌های CL ، ML ، WLS و MWL بر حسب تعداد زمان و مکان متفاوت اندازه‌گیری و نتایج در جدول ۲ آمده است. همانطوری که از جدول ۲ ملاحظه می‌شود بجز در مورد روش WLS که به دلیل دقت کمتر آماری مورد نظر ما نیست، کارایی زمانی روش MWL به ویژه در ابعاد بالای تعداد مشاهدات نسبت به روش‌های CL و ML به مراتب بهتر است.

۲.۴ شبیه‌سازی با مدل تفکیک‌ناپذیر

در ادامه یک مدل کاملاً متقارن تفکیک‌ناپذیر استفاده شده است. کاملاً متقارن بودن معادل این شرط است که $H(s, \omega)$ حقیقی مقدار است. در این شبیه‌سازی از تابع کواریانس کاملاً متقارن تفکیک‌ناپذیر مثال ۳.۲ استفاده خواهد شد. در این شبیه‌سازی پارامترهای تعامل فضایی-زمانی و همواری فضایی مقادیر ثابت $\beta = 0.5$ و $\alpha = 1$ در نظر گرفته شده‌اند. نتایج برازش بر اساس ۱۰۰۰ تکرار در جدول ۳ آمده است. مانند مثال قبل با افزایش n و یا MSE ، T برآوردها به کندی کاهش می‌یابد. با توجه به مقادیر MSE در جدول‌های ۱ و ۳ بهترین روش برای برآورد پارامترهای مدل، روش ML و پس از آن روش MWL مناسب است. همانطوری که قبلاً ذکر شد در حالتی که تعداد مشاهدات بزرگ است، روش ML دشوار و به صرفه نیست، بنابراین در این مقاله برای برآورد پارامترهای مدل برازش شده به داده‌های واقعی از روش MWL استفاده خواهد شد.

برای مدل تفکیک‌ناپذیر مثال ۳.۲ با $\beta = 0.5$ و $\alpha = 1$ زمان لازم برای یک تکرار روش‌های CL ، ML ، WLS و MWL بر حسب تعداد زمان و مکان متفاوت اندازه‌گیری و نتایج در جدول ۴ آمده است. همانطوری که از جدول ۴ ملاحظه می‌شود در این حالت نیز بجز روش WLS که به دلیل دقت کمتر آماری مورد نظر ما نیست، کارایی زمانی روش MWL به ویژه در ابعاد بالای تعداد مشاهدات نسبت به روش‌های CL و ML بهتر است.

جدول ۱: مقایسه MSE در داده‌های شبیه‌سازی شده برای مدل تفکیک‌پذیر مثال ۱.۲

$n = 20$		$n = 10$		روش برآورد	پارامترها
$T = 90$	$T = 60$	$T = 90$	$T = 60$		
MSE	MSE	MSE	MSE		
۰/۰۰۸	۰/۰۱۰	۰/۰۴۰	۰/۰۴۵	CL	اثر قطعه‌ای
$0.4e - 3$	۰/۰۰۶	۰/۰۱۱	۰/۰۱۸	ML	
۰/۰۵۰	۰/۱۰۰	۰/۱۲۴	۰/۱۴۰	WLS	
۰/۰۰۱	۰/۰۱۰	۰/۰۲۵	۰/۰۳۲	MWL	
۰/۰۰۸	۰/۰۱۰	۰/۰۶۱	۰/۰۶۴	CL	مقیاس فضایی
$0.4e - 1$	۰/۰۰۳	۰/۰۱۲	۰/۰۶۰	ML	
۰/۱۰۴	۰/۱۱۸	۰/۱۳۶	۰/۴۱۰	WLS	
۰/۰۰۷	۰/۰۵۷	۰/۰۵۳	۰/۰۶۴	MWL	
۰/۰۲۰	۰/۰۳۲	۰/۰۴۱	۰/۰۴۶	CL	مقیاس زمانی
۰/۰۰۷	۰/۰۱۶	۰/۰۲۵	۰/۰۳۶	ML	
۰/۰۷۴	۰/۱۹۷	۰/۱۴۳	۰/۱۸۵	WLS	
۰/۰۰۵	۰/۰۲۹	۰/۰۳۷	۰/۰۴۶	MWL	
۰/۰۰۴	۰/۰۱۸	۰/۰۲۸	۰/۰۲۴	CL	ازاره
$4e - 6$	۰/۰۰۹	۰/۰۱۹	۰/۰۱۵	ML	
۰/۰۸۴	۰/۱۰۰	۰/۱۴۰	۰/۱۳۸	WLS	
۰/۰۲۰	۰/۰۱۰	۰/۰۲۱	۰/۰۲۴	MWL	

جدول ۲: زمان لازم (برحسب ثانیه) برای تکرار روش‌های CL ، ML ، WLS و MWL در مدل تفکیک‌پذیر برحسبتعداد زمان و مکان متفاوت $n = 10, 20$ و $T = 60, 90$

MWL	WLS	ML	CL	n	T
۳/۳۷	۰/۳۱	۲۶۳/۸۵	۳۴/۹۱	۱۰	۶۰
۳/۹۷	۰/۶۹	۳۰۴/۸۰	۷۲/۰۸	۱۰	۹۰
۱۴/۸۶	۱/۰۸	۹۸۳/۵۲	۱۰۱/۲۸	۲۰	۶۰
۱۵/۵۳	۲/۶۴	۳۱۰۳/۹۲	۲۳۶/۶۸	۲۰	۹۰

۳.۴ تحلیل داده‌های واقعی

داده‌های مورد بررسی، مقادیر میانگین سرعت باد روزانه بر حسب متر بر ثانیه در ۸ ایستگاه هواشناسی استان زنجان مربوط به سال‌های ۲۰۰۳ تا ۲۰۱۷ است که توسط اداره هواشناسی استان زنجان اندازه‌گیری و ثبت شده‌اند. به دلیل وفور داده‌های گمشده در ایستگاه‌های مختلف، تنها از داده‌های ۵ ایستگاه که به طور کامل اطلاعات مربوط به سرعت باد طی سال‌های یاد شده را دارند، استفاده شده است. جدول ۵ طول و عرض جغرافیایی ایستگاه‌های مورد نظر را نشان می‌دهد که در آن طول و عرض جغرافیایی هر ایستگاه برحسب درجه، دقیقه و ثانیه، ارتفاع از سطح دریا برحسب متر، میانگین و انحراف معیار سرعت باد در هر ایستگاه بر حسب متر بر ثانیه ثبت شده است.

برای تحلیل فضایی-زمانی داده‌ها نخست تحلیل کاوشگرانه مشاهدات از نظر نرمال بودن، مانایی و همگنی در واریانس انجام می‌شود. نمودار بافت نگار داده‌ها در شکل ۱ سمت چپ، نشان‌دهنده عدم تقارن و انحراف از نرمال بودن توزیع

جدول ۳: مقایسه MSE در داده‌های شبیه‌سازی شده برای مدل تفکیک‌ناپذیر مثال ۳.۲ یا $\alpha = 1$ ، $\beta = 0.5$

$n = 20$		$n = 10$		روش برآورد	پارامترها
$T = 90$	$T = 60$	$T = 90$	$T = 60$		
MSE	MSE	MSE	MSE		
۰/۰۴۱	۰/۲۲۹	۰/۴۱۴	۰/۵۹۱	CL	اثر قطعه‌ای
۰/۰۰۵	۰/۱۰۷	۰/۲۵۳	۰/۴۶۸	ML	
۰/۱۳۱	۰/۳۶۲	۰/۶۳۸	۰/۷۴۱	WLS	
۰/۰۱۲	۰/۱۵۶	۰/۳۱۲	۰/۵۷۳	MWL	
۰/۰۴۳	۰/۱۶۷	۰/۲۸۹	۰/۲۸۹	CL	مقیاس فضایی
۰/۰۰۳	۰/۰۵۶	۰/۱۲۰	۰/۲۳۸	ML	
۰/۱۱۴	۰/۲۰۱	۰/۴۱۸	۰/۴۲۶	WLS	
۰/۰۱۲	۰/۱۴۳	۰/۲۱۷	۰/۲۵۴	MWL	
۰/۰۱۵	۰/۱۳۶	۰/۲۶۵	۰/۳۸۴	CL	مقیاس زمانی
۰/۰۰۹	۰/۰۸۷	۰/۲۵۱	۰/۳۲۸	ML	
۰/۲۶۱	۰/۳۲۵	۰/۵۹۳	۰/۵۷۸	WLS	
۰/۰۶۳	۰/۱۰۴	۰/۲۸۵	۰/۳۳۲	MWL	
۰/۰۰۸	۰/۰۷۴	۰/۱۲۷	۰/۱۹۳	CL	ازاره
۰/۰۰۱	۰/۰۰۸	۰/۰۶۸	۰/۱۵۲	ML	
۰/۱۱۶	۰/۱۸۱	۰/۲۴۲	۰/۳۶۱	WLS	
۰/۰۱۲	۰/۰۵۴	۰/۱۰۲	۰/۱۶۱	MWL	

جدول ۴: زمان لازم (برحسب ثانیه) برای تکرار روش‌های CL ، ML ، WLS و MWL در مدل تفکیک‌ناپذیر برحسب

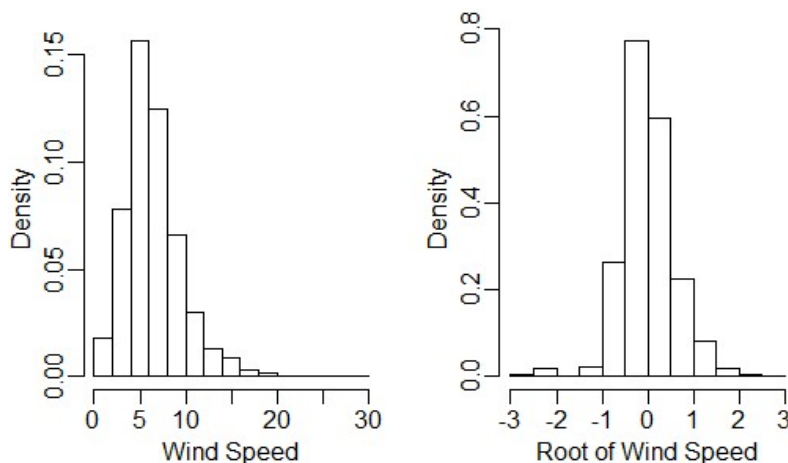
تعداد زمان و مکان متفاوت $n = 10, 20$ و $T = 60, 90$

MWL	WLS	ML	CL	n	T
۱۰/۸۸	۰/۳۶	۵۷/۶۴	۱۶/۷۹	۱۰	۶۰
۳۲/۹۸	۰/۸۰	۸۷/۷۰	۵۴/۵۰	۱۰	۹۰
۳۸/۸۸	۱/۲۲	۱۶۶/۵۶	۹۳/۴۴	۲۰	۶۰
۱۰۴/۹۲	۲/۹۴	۶۶۳/۵۸	۳۷۲/۱۱	۲۰	۹۰

جدول ۵: شماره، نام و موقعیت جغرافیایی ایستگاه‌ها

شماره	نام ایستگاه	طول و عرض جغرافیایی		ارتفاع از سطح دریا	میانگین سرعت باد	انحراف معیار سرعت باد
		عرض	طول			
۱	ماه‌نشان	$36^{\circ}44'N$	$47^{\circ}41'E$	۱۳۵۰	۶/۳۲۱	۲/۳۶۸
۲	زنجان	$36^{\circ}31'E$	$36^{\circ}39'N$	۱۶۶۳	۶/۷۴۸	۲/۸۱۳
۳	خرم‌دره	$49^{\circ}12'E$	$36^{\circ}11'N$	۱۵۷۵	۷/۳۳	۲/۸۱۸
۴	خدابنده	$48^{\circ}35'E$	$36^{\circ}08'N$	۱۹۹۷	۸/۷۰۲	۳/۷۳۰
۵	ابهر	$48^{\circ}56'E$	$36^{\circ}56'N$	۱۵۴۰	۵/۶۸۱	۳/۰۱۹

داده‌ها است. به منظور نرمال‌سازی داده‌ها از دو تبدیل ریشه دوم و لگاریتم استفاده شده است. اما تبدیل ریشه دوم داده‌ها که نمودار آن در شکل ۱ سمت راست رسم شده، موجب تقارن و به طور تقریبی نرمال شدن توزیع داده‌ها شده است. همچنین باید فرض همگنی واریانس روی فضا و زمان برای داده‌های اصلی بررسی شود. **ساهو و همکاران (۲۰۰۶)** و **هوانگ و همکاران (۲۰۰۷)** برای بررسی فرض همگنی واریانس در داده‌های فضایی-زمانی از نمودار انحراف استاندارد نمونه در مقابل میانگین نمونه استفاده کردند. با رسم این نمودار برای متوسط‌های زمانی در شکل ۲ سمت چپ، وجود روند و انحراف از فرض همگونی واریانس مشاهده می‌شود که خوشبختانه این روند و ناهمگونی در واریانس با بکار بردن تبدیل ریشه دوم داده‌های اصلی، همانطوری که در شکل ۲ سمت راست، ملاحظه می‌شود حذف شده است. همچنین همانطوری که در شکل ۳ سمت راست، ملاحظه می‌شود در داده‌های اصلی روند فصلی وجود دارد که این روند فصلی به کمک روش برازش رگرسیون ناپارامتری با استفاده از تابع *Lowess* در *R* حذف شده است. این تابع محاسبات هموارسازی *Lowess* را با استفاده از رگرسیون چندجمله‌ای وزنی موضعی انجام می‌دهد (**کلوند، ۱۹۸۱**). شکل ۳ سمت چپ، به وضوح حذف این اثر فصلی را نشان می‌دهد. نهایتاً مدل فضایی-زمانی ارائه شده روی باقیمانده‌های حاصل از حذف روند انجام گرفته است. برازش مدل‌های معرفی شده در بخش قبل به داده‌ها می‌تواند با محاسبه تابع کواریانس از طریق حل عددی



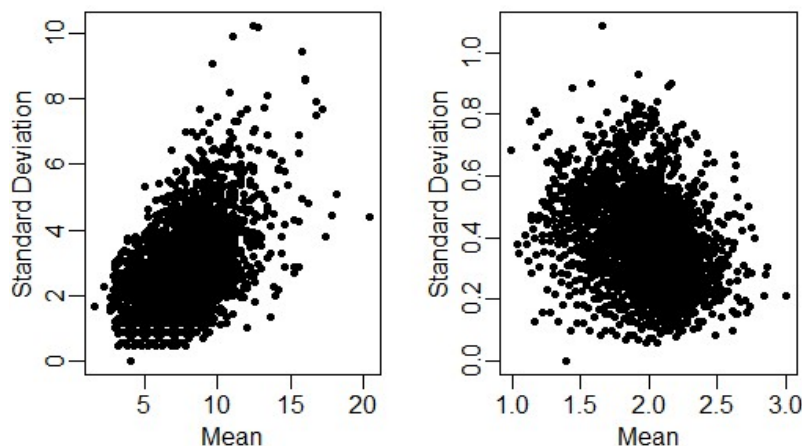
شکل ۱: هیستوگرام داده‌های اصلی و ریشه دوم داده‌ها

انتگرال رابطه (۱.۲) انجام شود و یا اینکه به داده‌ها به عنوان سری زمانی چند متغیره نگاه کرد و روش‌های طیفی را بکار برد. در این مقاله از برازش‌های طیفی استفاده و در روش طیفی برای برآورد پارامترها از درستنمایی چند متغیره وایتل استفاده شده است. مدل‌های معرفی شده در بخش قبل در زمان و فضا پیوسته هستند، در حالی که روش درستنمایی وایتل برای سری‌های زمانی گسسته بکار برده می‌شود. به منظور رفع این مشکل تصحیح هم‌اثرسازی^۱ زیر بکار رفته و برای مدل پیوسته $f(\omega)\mathbb{C}(\text{sd}(\omega))$ از تقریب

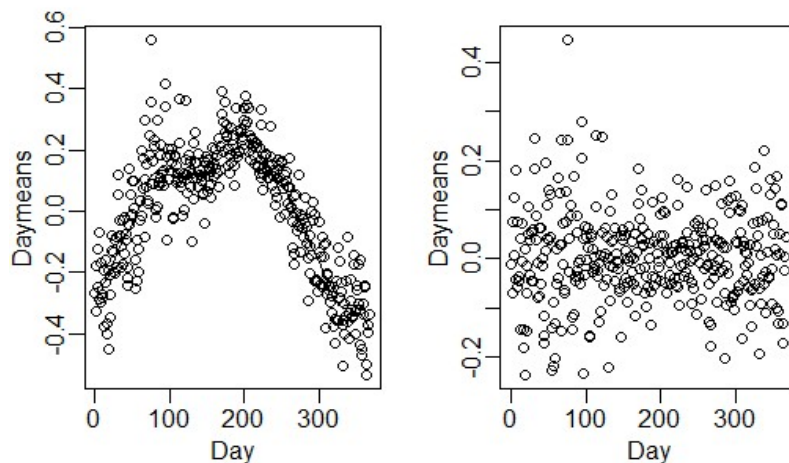
$$\sum_{j=-m}^m f(\omega + 2j\pi)\mathbb{C}(\text{sd}(\omega + 2j\pi)), \quad -\pi < \omega < \pi, \quad m = 50$$

استفاده شده است. نتایج برازش مدل‌های ذکر شده به داده‌ها و مقادیر لگاریتم درستنمایی وایتل و همچنین تقریبی از معیار اطلاع آکائیک (*AIC*) برای هریک از این مدل‌ها در جدول ۶ ثبت شده است. هرچند که معیار اطلاع آکائیک یکی از کاربردی‌ترین ابزارهای انتخاب مدل می‌باشد، ولی تحقیقات کمتری در مورد کارایی این ابزار در مدل‌سازی فضایی و

¹Aliasing Correction



شکل ۲: نمودار انحراف معیار نمونه در مقابل میانگین نمونه قبل (چپ) و بعد از تبدیل ریشه دوم (راست)



شکل ۳: حذف روند فصلی در داده‌ها

فضایی-زمانی انجام شده است. **هوتینگ و همکاران (۲۰۰۶)** و **همچنین لی و گاش (۲۰۰۹)** تأثیر همبستگی فضایی روی انتخاب مدل توسط معیار AIC را در آمار فضایی مورد بررسی قرار داده‌اند و نتایج شبیه‌سازی آنها نشان داد که معیار AIC با در نظر گرفتن همبستگی فضایی بهتر از حالتی است که ساختار فضایی در نظر گرفته نشده است. از طرفی آنها نشان داده‌اند که اگر اندازه نمونه بزرگ باشد، معیار AIC پیشنهادی آنها تقریباً برابر معیار AIC در حالتی که داده‌ها مستقل هستند می‌باشد. علاوه بر این **هوانگ و همکاران (۲۰۰۷)** نشان دادند که AIC هنوز هم می‌تواند معیار مناسبی برای انتخاب مدل بین مدل‌های فضایی-زمانی باشد. در این مقاله نیز به پیروی از مقالات (**هوتینگ و همکاران، ۲۰۰۶**) و (**هوانگ و همکاران، ۲۰۰۷**) به دلیل اینکه مجموع نمونه‌های فضایی و زمانی بسیار بزرگ است، برای سادگی از AIC بدون در نظر گرفتن همبستگی داده‌ها استفاده شده است و تعداد پارامترهای موثر تقریباً برابر تعداد پارامترهای مدل در نظر گرفته شده‌اند. با مقایسه مقادیر لگاریتم درستنمایی و ایتل چند متغیره و معیار اطلاع آکائیک (AIC) در جدول ۶ نتیجه می‌شود که در بین پنج مدل استفاده شده، مدل مربوط به مثال ۳.۲، برازنده‌ترین مدل برای داده‌های باد استان زنجان بوده است.

جدول ۶: مقایسه درست‌نمایی وایتل در مدل‌های برازش شده به داده‌های باد استان زنجان

مدل	مقدار ماکسیمم	اختلاف با مدل مثال ۳.۲	تعداد پارامترها	AIC
مثال ۱.۲	۲۵۵۲۵/۲	۷۰۰۴۵/۷	۴	-۵۱۰۴۲/۴
مثال ۲.۲	۲۵۵۴۹/۸	۷۰۰۲۱/۱	۴	-۵۱۰۹۱/۶
مثال ۳.۲	۹۵۵۷۰/۹	۰	۶	-۱۹۱۱۲۹/۸
مثال ۴.۲	۶۰۱۴۴/۲۵	۳۵۴۲۶/۷	۵	-۱۲۰۲۷۸/۵
مثال ۵.۲	۳۳۰۴۹	۶۲۵۲۱/۹	۴	-۶۶۰۹۰

با در نظر گرفتن اثر قطعه‌ای، مدل نیمه‌طیفی مثال ۳.۲ به صورت زیر در نظر گرفته شده است:

$$H(\mathbf{s}, \omega) = \begin{cases} \eta^2 \sigma^2 \times \frac{2a}{a^2 + \omega^2} & \|\mathbf{s}\| = 0 \\ \frac{\sigma^2}{(b \|\mathbf{s}\|)^\alpha + 1} \times \frac{2a((b \|\mathbf{s}\|)^\alpha + 1)^{-\frac{\beta}{\alpha}}}{a^2((b \|\mathbf{s}\|)^\alpha + 1)^{-\beta} + \omega^2} & \|\mathbf{s}\| \neq 0 \end{cases}$$

با استفاده از تابع درست‌نمایی وایتل چند متغیره در رابطه (۱.۳) برای برآورد پارامترهای $H(\mathbf{s}, \omega)$ یعنی $\theta = (\eta^2, \sigma^2, a, b, \alpha, \beta)'$ مقادیر

$$\hat{\eta} = 0.029, \hat{\sigma}^2 = 0.267, \hat{a} = 0.0012, \hat{b} = 1/121, \hat{\alpha} = 1/13, \hat{\beta} = 1/57$$

و مقدار لگاریتم درست‌نمایی وایتل مربوطه $MWL(\theta) = 95570.9$ بدست آمده است.

۵ بحث و نتیجه گیری

در این مقاله تحلیل نیمه‌طیفی داده‌های فضایی - زمانی توسط مدل‌های نیمه‌طیفی مورد بررسی و مطالعه قرار گرفت. یک مطالعه شبیه‌سازی به منظور مقایسه عملکرد روش نیمه‌طیفی با روش‌های کلاسیک با توجه به مدل‌های کوواریانس فضایی - زمانی تفکیک‌پذیر و تفکیک‌ناپذیر انجام شد. برای این منظور از چهار روش برآوردیابی ماکسیمم درست‌نمایی، درست‌نمایی ترکیبی، حداقل مربعات وزنی و روش نیمه‌طیفی استفاده و روش‌ها بر اساس میانگین مجذور خطاها با ۱۰۰۰ تکرار از میدان تصادفی گاوسی با میانگین صفر و واریانس یک مقایسه شده‌اند. نتایج شبیه‌سازی گویای عملکرد بهتر روش نیمه‌طیفی نسبت به دیگر روش‌ها هم از نظر محاسباتی و هم از نظر آماری است. در پایان نیز تحلیل نیمه‌طیفی داده‌های میانگین سرعت باد روزانه استان زنجان به عنوان کاربردی از مدل پیشنهادی ارائه شد.

مراجع

امیدی، م. و محمدزاده، م. (۱۳۹۲). تعیین ساختار همبستگی داده‌های فضایی با توابع مفصل، نشریه علوم دانشگاه خوارزمی، ۳، شماره ۳، ۷۹۷-۸۰۸.

Besag, J. (1974), *Spatial Interaction and the Statistical Analysis of Lattice Systems with Discussion*, Journal of the Royal Statistical Society, Series B, 192-236.

Bochner. S. (1955), *Harmonic Analysis and The Theory of Probability*. University of California Press, Berkely and Los Angeles.

Chiles, J. P. and Delfner, P. (2012), *Geostatistics: Modeling Spatial Uncertainty*. 2nd ed., Wiley, New York.

Cleveland, W. S. (1981), LOWESS: *A program for smoothing scatterplots by robust locally weighted regression*. The American Statistician 35, 54.

Cressie, N. (1985), *Fitting variogram models by weighted least squares*. Journal of the International Association for Mathematical Geology, 17(5), 563-586.

Cressie, N. (1993), *Statistics for Spatial Data*. Revised edn, New York: Wiley.

Cressie, N. and Huang, H. C. (1999), Classes of nonseparable, spatio-temporal stationary covariance functions. Journal of the American Statistical Association, 94, 1330-1340.

Cressie, N. and Wikle, C. K. (2013), *Statistics for spatio-temporal data*, John Wiley and Sons.

Gneiting, T. (2002), *Nonseparable stationary covariance functions for space-time data*. Journal of the American Statistical Association, 97, 590-600.

Hoeting, J. A., Davis, R. A., Merton, A. A., & Thompson, S. E. (2006), *Model selection for geostatistical models*. Ecological Applications, 16, 87-98.

Horrell, M. T. and Stein, M. L. (2017), *Half-spectral space-time covariance models*. Spat. Stat., 19, 90-100.

Huang, H. C., Martinez, F., Mateu, J and Montes, F. (2007), *Model comparison and selection for stationary space-time models*. Computational statistics and data analysis, 51, 4577-4596.

Journel, A. G. and Huijbregts, C. J. (1978), *Mining geostatistics*. Academic press.

Kammler, D. W. (2007), *A first course in Fourier analysis*. 2nd ed., Cambridge University Press, Cambridge.

Kent, J. T., Mohammadzadeh, M. and Mosammam, A. M. (2011), *The dimple in Gneiting's spatial-temporal covariance model*. Biometrika, 98, 489-494.

Lee, H., and Ghosh, S. K. (2009), *Performance of information criteria for spatial models*. Journal of statistical computation and simulation, 79(1), 93-106.

Lindsay, B. G. (1988), *Composite likelihood methods*. Contemporary Mathematics, 80, 221-39.

Mardia, K. V. and Marshall, R. J. (1984), *Maximum likelihood estimation of models for residual covariance in spatial regression*. Biometrika, 71, 135-146.

Mosammam, A. M. and Kent, J. T. (2016), *Estimation and testing for covariance-spectral spatial-*

temporal models. Environ. Ecol. Stat., 23, 43-64.

Omidi, M. and Mohammadzadeh, M., (2015), *A New Method to Build Spatio-Temporal Covariance Functions: Analysis of Ozone Data*. Statistical Papers, 57, 689–703.

Rodriguez-Iturbe, I. and Meji'a, J. M. (1974), *The design of rainfall networks in time and space*. WaterResources Research, 10, 713-728.

Sahu, S. K. G., Gelfand, A. E., Holland, D. M. and Mardia, K. (2006), *Spatio-Temporal modeling of fine particulate matter*. Journal of Agricultural, Biological and Environmental Statistics, 11, 61-86.

Shumway, R. H. and Stoffer, D. S. (2017), *Time Series Analysis and its Applications with R Examples*. 4th ed., Springer Texts in Statistics, Springer, Cham.

Stein, M. L. (2005), *Statistical methods for regular monitoring data*. J. R. Stat. Soc. Ser. B Stat.Methodol., 67, 667-687.

Stein, M. L. (2011), *When does the screening effect not hold?*. Ann. Statist., 39, 2795-2819.

Whittle, P. (1954), *On stationary processes in the plane*. Biometrika, 41, 434-449.



مقایسه دو تابع درستنمایی مرکب در برآورد مدل های آمیخته خطی تعمیم یافته فضایی

لیلا صالحی^۱، محسن محمدزاده
گروه آمار، دانشگاه تربیت مدرس

چکیده: هنگامی که با متغیرهای همبسته فضایی ناگوسی سر و کار داریم، مدل های آمیخته خطی تعمیم یافته فضایی از انعطاف لازم برای مدل بندی داده های گوناگون برخوردارند. اگر حجم داده ها بزرگ باشد، معمولاً روش های ماکسیمم درستنمایی در برآورد پارامترهای مدل با محاسبات سنگین و سرعت پایین همراه هستند. برای رفع این مشکل، توابع درستنمایی مرکب که از حاصل ضرب درستنمایی زیر مجموعه های مشاهدات به دست می آیند، به کار گرفته می شوند. در این مقاله برآورد پارامترهای مدل های آمیخته خطی تعمیم یافته فضایی با استفاده از تابع درستنمایی زوجی بررسی می شود. سپس تابع درستنمایی زوجی موزون و تابع درستنمایی زوجی توانیده برای برآورد پارامترهای این مدل ها را مورد استفاده قرار داده و در مطالعه ای شبیه سازی دقت برآوردها بر اساس این توابع درستنمایی مورد ارزیابی و مقایسه قرار می گیرد.

واژه های کلیدی: مدل های آمیخته خطی تعمیم یافته فضایی، تابع درستنمایی مرکب، درستنمایی زوجی، درستنمایی زوجی توانیده، درستنمایی زوجی موزون.
کد موضوع بندی ریاضی (۲۰۱۰): 91B72, 62M30, 62H11

۱ مقدمه

مدل های خطی تعمیم یافته برای اولین بار توسط **نلدر و وودرین (۱۹۷۲)** معرفی شدند و **مک-کالا (۱۹۸۹)** برای مدل بندی متغیرهای پاسخ گسسته از این مدل ها استفاده کردند. در این مدل ها با فرض استقلال مشاهدات، با استفاده از یک تابع پیوند بین میانگین مشاهدات و متغیرهای کمکی ارتباط خطی برقرار می شود. در حالتی که بین مشاهدات همبستگی وجود دارد، تعمیمی از مدل های مذکور به نام مدل های آمیخته خطی تعمیم یافته^۱ (GLM) استفاده می شود. در این مدل ها فرض استقلال مشاهدات به استقلال شرطی تعدیل و همبستگی بین آنها با اضافه کردن اثرات تصادفی از طریق متغیرهای پنهان^۲

^۱Generalized Linear Mixed

^۲Latent Variables

^۱ نام و ایمیل ارائه دهنده مقاله: لیلا صالحی، l.salehi@modares.ac.ir

به مدل در نظر گرفته می‌شود. در حالتی که ساختار همبستگی پاسخ‌های فضایی در قالب یک مدل مشخص باشد، با در نظر گرفتن یک میدان تصادفی تحلیل فضایی و پیشگویی مقادیر نامعلوم در موقعیت‌های مشخص با روش کریگیدن^۳ امکان‌پذیر است. اما در مواردی که ساختار همبستگی پاسخ‌های فضایی نامشخص باشد، می‌توان از مدل‌های آمیخته خطی تعمیم‌یافته فضایی^۴ (SGLM) استفاده نمود. چون در مدل‌های SGLM تابع درستنمایی برخلاف مدل‌های خطی به دلیل ناگوسی بودن متغیر پاسخ فرم بسته‌ای ندارد، برآورد پارامترها به روش ماکسیمم درستنمایی میسر نیست. لذا در اکثر مقالات با پذیرش فرض نرمال بودن متغیرهای پنهان به ارائه راه حلی برای برآورد پارامترهای مدل و متغیرهای پنهان با ماکسیمم کردن توابع درستنمایی، شبه‌درستنمایی توانیده^۵ یا درستنمایی سلسله مراتبی به روش‌های عددی پرداخته شده است. از جمله مک-کولاج (۱۹۹۷) با به کار بردن روش‌های عددی مثل ماکسیمم سازی امید ریاضی مونت کارلویی^۶ (MCEM)، برآورد پارامترهای مدل‌های GLM با اثرات تصادفی غیرفضایی را به دست آورد. محمدزاده و حسینی (۲۰۱۱) با فرض چوله بسته بودن متغیرهای پنهان الگوریتم گرادیانت EM مونت کارلو را برای برآورد ماکسیمم درستنمایی پارامترهای مدل پیشنهاد کردند. باغیشی و محمدزاده (۲۰۱۱) و ترابی (۲۰۱۵) روش دیگری بر اساس الگوریتم همسان‌سازی داده‌ها ارائه کردند. اخیراً محققان روش‌های تقریبی در قالب رهیافت بیزی و بسامدی را مورد توجه قرار داده‌اند که نسبت به روش‌های مبتنی بر الگوریتم‌های شبیه‌سازی، بار محاسباتی کمتر و سرعت بیشتری دارند. راثو و مارتینو (۲۰۰۷) استنباط بیزی تقریبی برای فرایندهای تصادفی مارکوفی گاوسی معرفی کردند. ایدسویک و همکاران (۲۰۱۲) روش مشابه را برای مدل‌های SGLM به کار بردند.

در مطالعات اخیر روش‌های تقریبی دیگری مورد توجه قرار گرفته‌اند که مبتنی بر درستنمایی کامل مشاهدات نیستند. مزیت این روش‌ها در مقایسه با روش‌های درستنمایی تقریبی، عدم نیاز به مدل‌بندی توأم تمام مشاهدات است. از جمله روش‌های شبه‌درستنمایی دودرین (۱۹۷۴) و مک-کالا (۱۹۸۳) و روش‌های مبتنی بر مشخصات جزئی درستنمایی کامل که درستنمایی مرکب^۷ نامیده می‌شود. برای اولین بار هگرتی و له (۱۹۹۸) تابع درستنمایی مرکب زوجی^۸ را برای مشاهدات دودویی با تابع پیوند پروبیت به کار بردند. وارین و همکاران (۲۰۰۵) روش درستنمایی مرکب زوجی را برای مدل‌های SGLM استفاده کردند. آن‌ها برای ماکسیمم کردن تابع درستنمایی زوجی، یک الگوریتم جدید از نوع EM را که از مربع‌سازی عددی استفاده می‌کند، معرفی کردند. بویلاکوا و همکاران (۲۰۱۰) برای داده‌های فضایی-زمانی از تابع درستنمایی موزون استفاده کردند و نشان دادند که برآورد بر اساس این تابع، تقریب مناسبی از برآورد ماکسیمم درستنمایی است. جو ولی (۲۰۰۹) نشان دادند که موزون کردن تابع درستنمایی موجب افزایش کارایی نسبی مجانبی در برآورد در داده‌های دسته‌بندی شده می‌شود.

در این مقاله تابع درستنمایی زوجی برای مدل‌های SGLM به کار گرفته می‌شود. سپس با استفاده از یک تابع وزن، تابع درستنمایی زوجی موزون را تشکیل داده و برای برآورد پارامترهای این مدل‌ها مورد استفاده قرار می‌گیرد. همچنین برای افزایش دقت در برآورد پارامترهای مدل، تابع درستنمایی زوجی توانیده به کار گرفته می‌شود و در مطالعه‌ای شبیه‌سازی، دقت برآورد پارامترهای مدل بر اساس توابع درستنمایی زوجی، درستنمایی زوجی موزون و درستنمایی زوجی توانیده بر اساس ملاک میانگین توان دوم خطا (MSE) مورد ارزیابی و مقایسه قرار می‌گیرد.

³Kriging

⁴Spatial Generalized Linear Mixed

⁵Penalized Quasi Likelihood

⁶Monte Carlo Expectation Maximization

⁷Composite Likelihood

⁸Pairwise Composite Likelihood

۲ مدل‌های آمیخته خطی تعمیم یافته فضایی

فرض کنید $Y(s)$ متغیر پاسخ فضایی گسسته و $Z_1(s), \dots, Z_p(s)$ متغیرهای تبیینی قابل مشاهده و $\{X(s), s \in \mathbb{R}^2\}$ یک میدان تصادفی فضایی پنهان باشد، طوری که $X(s)$ نشان دهنده اثر تصادفی در موقعیت s است. **دیگل و همکاران (۱۹۹۸)** مدل‌های SGLM را به صورت زیر تعریف کردند.

الف- $\{X(s), s \in \mathbb{R}^2\}$ یک میدان تصادفی مانای گاوسی است، به طوری که برای هر s ، $E[X(s)] = 0$ و برای هر $h \in \mathbb{R}$ ، $Cov(X(s+h), X(s)) = C(h; \theta)$ ، که در آن $\theta \in \mathbb{R}^k$ بردار پارامتر همبستگی است.

ب- $\{Y(s), s \in \mathbb{R}^2\}$ به شرط $\{X(s), s \in \mathbb{R}^2\}$ ، مجموعه‌ای از متغیرهای تصادفی مستقل است و توزیع هر $Y(s)$ با میانگین شرطی $E[Y(s)|X(s)]$ مشخص می‌شود.

ج- برای هر تابع پیوند g و پارامترهای رگرسیونی $\beta_1, \dots, \beta_p$ داریم $\beta_1, \dots, \beta_p$ داریم $g\{E[Y(s)|X(s)]\} = \sum_{j=1}^p Z_j(s)\beta_j + X(s)$ د- توزیع شرطی $[Y(s)|X(s)]$ عضو خانواده نمایی است.

۱.۲ تابع درستنمایی مرکب زوجی موزون

تابع درستنمایی مرکب از ضرب مجموعه‌ای از مولفه‌های درستنمایی به دست می‌آید، که هر مولفه درستنمایی زیرمجموعه‌ای از مشاهدات را بیان می‌کند. دو انگیزه اصلی برای استفاده از درستنمایی مرکب وجود دارد: انگیزه اول کاهش محاسبات در مدل‌بندی توزیع توام یک بردار تصادفی با بعد بالا است. انگیزه دوم استواری تحت نامعین بودن توزیع‌های با بعد بالاست. چون درستنمایی مرکب تنها به فرض‌های مربوط به چگالی‌های حاشیه‌ای و شرطی با بعد پایین نیاز دارد و آگاهی از جزئیات توزیع توام ضروری نخواهد بود. استنباط بر اساس تابع درستنمایی برای داده‌های حجیم با انتگرال‌گیری و معکوس ماتریس‌های بعد بالا همراه است که ممکن است حل آن‌ها برای رایانه‌های قدرتمند نیز مشکل باشد. با استفاده از تابع درستنمایی مرکب موزون می‌توان این انتگرال‌های چندگانه را به جمع انتگرال‌های با بعد پایین تبدیل نمود. بردار تصادفی m بعدی $Y = (y_1, \dots, y_m)$ را با تابع چگالی احتمال $f(y; \theta)$ برای پارامتر p بعدی $\theta = (\theta_1, \dots, \theta_m) \in \Theta$ در نظر بگیرد. فرض کنید $\{A_1, \dots, A_K\}$ مجموعه‌ای از پیشامدهای شرطی یا حاشیه‌ای با تابع درستنمایی $L_k(\theta; y) \propto f(y \in A_k; \theta)$ باشد. آنگاه تابع درستنمایی مرکب برای پارامتر θ به صورت

$$L_C(\theta; y) = \prod_{k=1}^K L_k(\theta, y)^{w_k}$$

تعریف می‌شود، که در آن w_k ها وزن‌هایی غیر منفی هستند (**لیندسی، ۱۹۹۸**). یک نمونه از توابع درستنمایی مرکب، تابع درستنمایی مرکب زوجی موزون است. **بویلاکوا و همکاران (۲۰۱۰)** تابع درستنمایی موزون را برای داده‌های فضایی-زمانی به کار بردند و نشان دادند برآوردهای درستنمایی مرکب موزون سازگار و به طور مجانبی دارای توزیع گاوسی با واریانس برابر معکوس اطلاع گودمب می‌باشند. آنها همچنین با مطالعات شبیه‌سازی نشان دادند که این برآورد تقریب مناسبی از برآورد ماکسیمم درستنمایی است و نسبت به برآوردهای ماکسیمم درستنمایی و درستنمایی مرکب دارای محاسبات کمتری است. **جو و لی (۲۰۰۹)** تابع درستنمایی مرکب موزون را برای داده‌های دسته‌بندی شده با تعداد دسته‌های متغیر به کار بردند و معیار کارایی نسبی مجانبی^۹ را برای وزن‌های مختلف بررسی کردند و نشان دادند که موزون کردن تابع درستنمایی

^۹ Asymptotic Relative Efficiencies

مرکب موجب افزایش کارایی نسبی مجانبی می‌شود. تابع درست‌نمایی مرکب زوجی موزون به صورت

$$L_W(\theta; y) = \prod_{r=1}^{m-1} \prod_{s=r+1}^m f(y_r, y_s; \theta)^{w_{rs}} \quad (1.2)$$

است. بنابراین تابع لگاریتم درست‌نمایی زوجی موزون به صورت

$$\ell_w(\theta; y) = \sum_r \sum_{s>r} w_{rs} \log f(y_r, y_s; \theta) \quad (2.2)$$

خواهد بود، که در اینجا برای مدل‌های SGLM از تابع

$$w_{i,j} = \begin{cases} 1 & \|s_i - s_j\| \leq d_s \\ 0 & \text{در غیر این صورت} \end{cases} \quad (3.2)$$

برای موزون کردن تابع درست‌نمایی مرکب زوجی استفاده می‌شود. در اینصورت برآورد ماکسیمم درست‌نمایی مرکب زوجی موزون θ که تابع (۲.۲) را ماکسیمم می‌کند، تحت شرایط نظم، برابر جواب منحصر به فرد معادله $u(\theta; y) = \nabla_{\theta} \ell_w = 0$ است. تابع درست‌نمایی زوجی برای مدل‌های SGLM به صورت

$$L_W(\eta; y) = \prod_{(i,j) \in \chi} L(\eta|y_i, y_j) \propto \prod_{(i,j) \in \chi} \int \int f(y_i|x_i) f(y_j|x_j) f(x_i, x_j|\eta) dx_i dx_j \quad (4.2)$$

است، که در آن χ زیرمجموعه همسایگی‌های زوجی (y_i, y_j) است.

۲.۲ تابع درست‌نمایی زوجی تاوانیده

روش ماکسیمم درست‌نمایی در برآورد پارامترها در مسائل گوناگون، گاهی با مسئله بیش برآزشی، دقت پایین یا واریانس زیاد برآوردگرها همراه است. تاوانیدن تابع درست‌نمایی، یک راه حل برای اصلاح چنین رفتارهایی در برآوردگرها به روش ماکسیمم درست‌نمایی معمولی است (آزالینی و واله، ۲۰۱۳) که روش ماکسیمم درست‌نمایی تاوانیده نامیده می‌شود. استفاده از تقریب تابع درست‌نمایی به جای درست‌نمایی معمولی به دلیل ماهیت تقریبی آن، می‌تواند بر کاهش دقت برآوردگرها اثر بگذارد. برای داده‌های فضایی حجیم که به دست آوردن تابع درست‌نمایی به روش تحلیلی غیرممکن است، می‌توان از تابع درست‌نمایی مرکب برای برآورد پارامترها استفاده کرد. تابع درست‌نمایی مرکب یک تقریب از تابع درست‌نمایی مشاهدات است، لذا ممکن است برآوردکننده‌ها بر اساس این تابع دقت پایینی داشته باشند. برای افزایش دقت برآورد در مدل‌های SGLM بر اساس این تابع، می‌توان از تاوانیدن تابع درست‌نمایی مرکب استفاده کرد. در این روش، به لگاریتم درست‌نمایی مشاهدات یک تابع تاوان اضافه می‌شود. برای مدل‌های SGLM، تابع درست‌نمایی زوجی تاوانیده به صورت

$$\ell(\eta; y) = \sum_r \sum_{s>r} \log f(y_r, y_s; \eta) - \lambda J(\eta)$$

است، که در آن λ پارامتر همواری و $J(\eta)$ تابع تاوان است که به صورت‌های گوناگون انتخاب می‌شود. برای مثال فرث (۱۹۹۳) تابع تاوان به صورت $\lambda \frac{1}{2} \log |I(\eta)|$ را برای خانواده نمایی به کار برد، که در آن $I(\eta)$ پیشین جفریز پارامتر η است. تیشیرانی (۱۹۹۶) تابع تاوان لاسو (LASSO) به صورت $J_{\lambda}(\eta) = \lambda \eta$ را برای برآورد در مدل‌های خطی ارائه کرد. در این مقاله از تابع تاوان $2\eta\eta^T$ استفاده شده است (گرین، ۱۹۹۰)، که در آن $\eta^T = (\beta_0, \beta_1, \sigma^2, \varphi)^T$ بردار پارامترهای مدل است.

۳.۲ الگوریتم EM برای ماکسیم سازی تابع درستنمایی زوجی

حسینی و محمدزاده (۱۳۹۲) الگوریتم گرادیانت ماکسیم سازی امیدریاضی زوجی تقریبی^{۱۰} (APEMG) را برای مدل های SGLM با متغیر پنهان فضایی ارائه کردند. این الگوریتم بر اساس تقریب تابع چگالی شرطی متغیر پنهان ارائه شده است. ایدسویک و همکاران (۲۰۱۲) نشان دادند که اگر X دارای توزیع نرمال و توزیع Y متعلق به خانواده نمایی باشد، با خطی کردن $f(y|x)$ حول یک مقدار ثابت x^* تابع چگالی یا جرم احتمال $f(x|y, \eta)$ را می توان با توزیع نرمال به صورت

$$\hat{f}(x|y, \eta) \approx N_n(\hat{\mu}_{x|y, \eta}(y, x^*), \hat{\Sigma}_{x|y, \eta}(x^*)) \quad (5.2)$$

تقریب زد، که در آن

$$\hat{\mu}_{x|y, \eta}(y, x^*) = H\beta + \Sigma_{\theta} R^{-1} (z(y, x^*) - H\beta),$$

$z_i(x_i^*, y_i) = [y_i - b'(x_i^*) + x_i^* b''(x_i^*)] / b''(x_i^*)$ ، $P_i i = 1/b''(x_i)$ با عناصر $n \times n$ قطری $P, R = \Sigma_{\theta} + P$ و $\hat{\Sigma}_{x|y, \eta}(x^*) = \Sigma_{\theta} - \Sigma_{\theta} R^{-1} \Sigma_{\theta}$ از رابطه (۵.۲) طبق خاصیت بسته بودن حاشیه سازی توزیع نرمال می توان نوشت $(x_{ij}|y_{ij}, \eta) \approx N_2(\hat{\mu}_{ij}, \hat{\Sigma}_{ij})$ ، که در آن $x_{ij} = (x_i, x_j)'$ ، $y_{ij} = (y_i, y_j)'$ و $\hat{\mu}_{ij}$ برداری دو بعدی با مولفه های برگرفته از سطرهای i و j از بردار $\hat{\mu}_{x|y}$ است. همچنین یک ماتریس 2×2 از سطر i و ستون j از ماتریس کواریانس $\hat{\Sigma}_{x|y}$ است. در این مقاله، برای ماکسیم کردن تابع درستنمایی زوجی موزون از الگوریتم APEM استفاده خواهد شد. گام های ۱ تا ۳ این الگوریتم همانند الگوریتم APEMG است و در گام ۴، برآورد پارامترها از ماکسیم کردن تابع

$$Q(\eta|\eta^{(m)}) = \sum_{(i,j) \in \chi} w_{ij} \int \int \log\{f(x_i, x_j, y_i, y_j; \eta)\} f(x_i, x_j|y_i, y_j; \eta^{(m)}) dx_i dx_j. \quad (6.2)$$

به دست می آید. برای محاسبه انتگرال دوگانه فوق از الگوریتم MCMC و تقریب تابع توزیع شرطی (۵.۲) استفاده می شود.

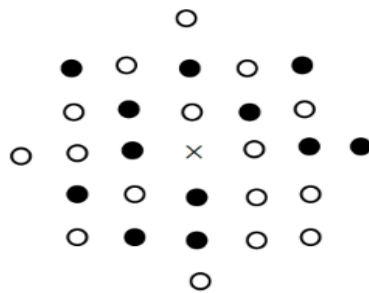
۳ مطالعه شبیه سازی

در این بخش تاثیر موزون کردن تابع درستنمایی زوجی بر دقت برآورد پارامترهای مدل SGLM در یک مطالعه شبیه سازی بررسی می شود. برای این منظور، داده ها از یک شبکه منظم 10×10 با گره های $\{(\ell, k), \ell, k = 1, \dots, 10\}$ تولید شده اند. برای تولید متغیرهای پنهان فضایی X از توزیع نرمال $N_{100}(\beta_1 z, \Sigma_{\theta})$ ، تابع کواریانس نمایی همسان گرد $C(h) = \sigma^2 \exp(-h/\varphi)$ و مقادیر $\sigma^2 = 1$ ، $\beta_1 = -1$ و $\beta_0 = 0.5$ و $\varphi = 2$ در نظر گرفته شده اند. متغیر تبیینی در هر موقعیت $s = (\ell, k)$ به صورت $Z_{\ell, k} = \log(1 + \ell)$ در نظر گرفته شده است. متغیرهای پاسخ $Y_{\ell, k}$ نیز با شرطی کردن روی متغیرهای پنهان فضایی از توزیع $Bin(100, \exp(x_{\ell, k}) / [1 + \exp(x_{\ell, k})])$ تولید شده اند. پارامترها با الگوریتم APEM برآورد شده اند. برای هر مشاهده ۱۲ زوج همسایه به صورت تصادفی بدون جایگذاری با همسایگی به شعاع ۳ انتخاب شده است. در شکل ۳ برآورد پارامترها با استفاده از تابع درستنمایی زوجی و درستنمایی زوجی وزنی انجام شده و دقت برآورد بر اساس ملاک MSE مورد مقایسه قرار گرفته است. برای به دست آوردن مقدار بهینه تابع وزن، مقادیر مختلف d_s در تابع (۲.۲) مورد بررسی قرار گرفته شده است. نتایج در جدول ۳ و نمودار ۲ نشان می دهد که برای برآورد پارامتر β_0 ، مقدار MSE تابع درستنمایی زوجی موزون برای همه مقادیر مختلف فاصله d_s از تابع درستنمایی زوجی معمولی کمتر است. در برآورد پارامتر β_1 ، تابع درستنمایی زوجی موزون با مقدار فاصله $d_s = 2/5$ دارای MSE کمتری است. برای برآورد پارامتر σ^2 تابع درستنمایی زوجی موزون با مقدار فاصله $d_s = 2/4$ بهترین عملکرد را دارد و تابع درستنمایی زوجی

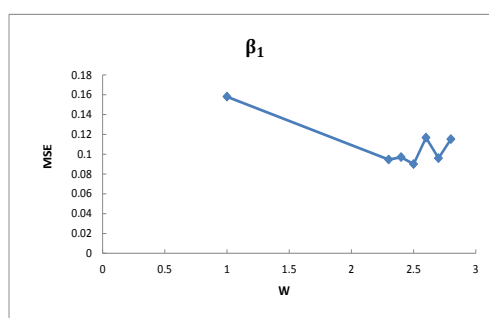
¹⁰ Approximate Pairwise EMG

جدول ۱: برآورد پارامترهای مدل SGLM بر اساس تابع درستنمایی زوجی و درستنمایی زوجی وزنی

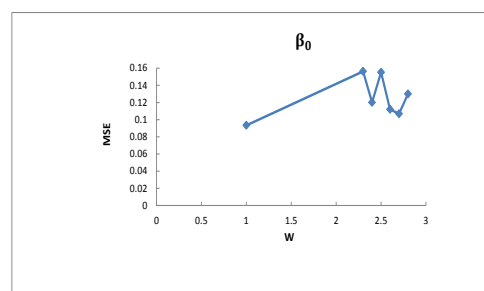
تابع درستنمایی	تابع وزن	پارامتر	برآورد	MSE	SE
معمولی	۱	β_0	-۰/۸۶۷۸	۰/۱۵۸۰	۰/۰۳۷۷
		β_1	۰/۵۶۵۳	۰/۰۹۳۴	۰/۰۳۰۵
		σ^2	۰/۷۱۲۵	۰/۱۱۳۷	۰/۰۱۷۷
		φ	۱/۶۶۶۶	۳/۴۱۱۱	۰/۱۸۲۶
۲/۳	۲/۳	β_0	-۰/۸۶۶۱	۰/۱۵۶۳	۰/۰۳۷۴
		β_1	۰/۴۶۳۸	۰/۰۹۴۶	۰/۰۳۰۷
		σ^2	۰/۷۱۵۲	۰/۱۱۱۶	۰/۰۱۷۶
		φ	۲/۰۰۷۸	۶/۸۱۱۴	۰/۳۷۵۶
۲/۴	۲/۴	β_0	-۰/۹۰۸۹	۰/۱۲۰۲	۰/۰۳۳۶
		β_1	۰/۴۷۳۶	۰/۰۹۷۰	۰/۰۳۱۲
		σ^2	۰/۷۴۹۹	۰/۰۹۷۰	۰/۰۱۸۷
		φ	۲/۳۰۵۷	۱۴/۰۵۶۴	۰/۳۷۵۶
موزون	۲/۵	β_0	-۰/۸۸۸۱	۰/۱۵۵۲	۰/۰۳۸۰
		β_1	۰/۴۶۹۰	۰/۰۹۰۰	۰/۰۳۰۰
		σ^2	۰/۶۹۸۹	۰/۱۲۴۲	۰/۰۱۸۴
		φ	۲/۱۶۱۰	۱۳/۴۱۴۳	۰/۳۶۷۷
۲/۶	۲/۶	β_0	-۰/۹۲۱۹	۰/۱۱۲۰	۰/۰۳۲۷
		β_1	۰/۵۱۸۹	۰/۱۱۶۷	۰/۰۳۴۳
		σ^2	۰/۷۱۰۵	۰/۱۱۱۱	۰/۰۱۶۶
		φ	۱/۶۸۵۳	۳/۴۸۱۹	۰/۱۸۴۹
۲/۷	۲/۷	β_0	-۰/۹۱۵۳	۰/۱۰۶۹	۰/۰۳۱۷
		β_1	۰/۴۷۹۱	۰/۰۹۶۰	۰/۰۳۱۱
		σ^2	۰/۶۸۴۵	۰/۱۴۳۶	۰/۰۲۱۱
		φ	۱/۹۵۸۵	۲/۹۶۸۴	۰/۱۷۳۱
۲/۸	۲/۸	β_0	-۰/۹۳۴۰	۰/۱۳۰۰	۰/۰۳۵۶
		β_1	۰/۴۸۷۶	۰/۱۱۵۲	۰/۰۳۰۰
		σ^2	۰/۷۲۵۴	۰/۱۱۲۹	۰/۰۱۹۴
		φ	۱/۹۹۳۷	۱۸/۴۹۸۳	۰/۴۳۲۳



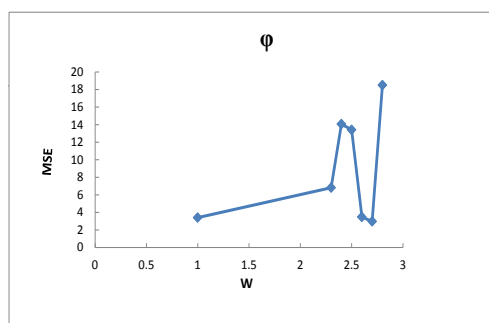
شکل ۱: نمونه‌گیری زوجی در همسایگی به شعاع ۳، موقعیت‌های با علامت ضربدر و دایره‌ای توپر یک زوج مرتب هستند



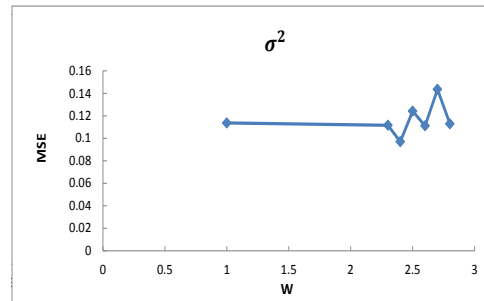
(ب)



(الف)



(د)



(ج)

شکل ۲: نمودار MSE بر اساس توابع درستنمایی معمولی و موزون برای پارامترهای مدل

موزون یا مقدار فاصله $d_s = 2/V$ در بین سایر توابع عملکرد بهتری در برآورد پارامتر همبستگی φ دارد. بنابراین استفاده از تابع درستنمایی زوجی موزون در برآورد پارامترهای مدل‌های آمیخته خطی تعمیم‌یافته فضایی دارای دقت بالاتری نسبت به تابع درستنمایی زوجی معمولی است.

در ادامه، برای بررسی تاثیر توانیدن تابع درستنمایی زوجی بر روی دقت برآورد پارامترها بر اساس معیار MSE، تابع درستنمایی زوجی توانیده مورد مطالعه قرار گرفته شده است. برای برآورد پارامترها از الگوریتم APEM استفاده شده که در مرحله M آن تابع $Q(\eta|\eta^m) - \lambda J(\eta)$ ماکسیمم می‌شود. همان‌طور که در جدول ۳ ملاحظه می‌شود، مقدار کمیت‌های MSE و SE در برآورد پارامترهای مدل با استفاده از تابع درستنمایی زوجی توانیده کمتر از مقدار این کمیت‌ها در تابع

جدول ۲: برآورد پارامترهای مدل با SGLM توابع درستنمایی زوجی و زوجی تاوانیده

تابع درستنمایی	پارامتر	برآورد	MSE	SE
زوجی تاوانیده	β_0	-۰/۸۹۰۸	۰/۱۳۷۳	۰/۰۳۵۵
	β_1	۰/۳۹۹۹	۰/۱۳۷۹	۰/۰۲۷۱
	σ^2	۰/۷۳۰۲	۰/۱۰۱۰	۰/۰۱۶۹
	φ	۱/۸۹۴۷	۲/۲۵۰۹	۰/۱۵۰۴
زوجی	β_0	-۰/۸۶۷۸	۰/۱۵۸۰	۰/۰۳۷۷
	β_1	۰/۵۶۵۳	۰/۰۹۳۴	۰/۰۳۰۵
	σ^2	۰/۷۱۲۵	۰/۱۱۳۷	۰/۰۱۷۷
	φ	۱/۶۶۶۶	۳/۴۱۱۱	۰/۱۸۲۶

درستنمایی زوجی است. برای مثال در برآورد پارامتر β_0 بر اساس تابع درستنمایی زوجی تاوانیده، مقدار MSE و SE به ترتیب ۰/۱۳۷۳ و ۰/۰۳۵۵ هستند در حالی که مقدار این کمیت‌ها در برآورد با استفاده از تابع درستنمایی زوجی، به ترتیب ۰/۱۵۸۰ و ۰/۰۳۷۷ هستند. در برآورد پارامترهای σ^2 و φ نیز استفاده از تابع درستنمایی تاوانیده دقت بالاتری نسبت به درستنمایی زوجی فراهم می‌کند. بنابراین می‌توان گفت که تاوانیدن تابع درستنمایی زوجی تا حدود زیادی توانسته دقت برآورد پارامترها را افزایش دهد.

بحث و نتیجه‌گیری

در این مقاله، تابع درستنمایی زوجی، درستنمایی زوجی موزون و درستنمایی زوجی تاوانیده برای مدل‌های آمیخته خطی تعمیم‌یافته فضایی به کار گرفته شد. برای ماکسیم‌سازی این توابع، الگوریتم APEM مورد استفاده قرار گرفت. بر اساس مشاهدات شبیه‌سازی شده، دقت برآورد پارامترها بر اساس این دو تابع و بکارگیری ملاک MSE مقایسه شد. نتایج شبیه‌سازی نشان دادند که تابع درستنمایی زوجی موزون و تابع درستنمایی زوجی تاوانیده در برآورد پارامترهای مدل آمیخته خطی تعمیم‌یافته فضایی، دارای دقت بالاتری نسبت به تابع درستنمایی زوجی هستند. بعلاوه تابع درستنمایی زوجی تاوانیده در برآورد پارامترهای همبستگی و تابع درستنمایی زوجی موزون در برآورد پارامترهای رگرسیونی دارای دقت بالاتری نسبت به سایر توابع بودند. در مقایسه با تابع درستنمایی زوجی موزون با توجه مقادیر MSE دو جدول، تابع درستنمایی زوجی تاوانیده در برآورد دو پارامتر σ^2 و φ دارای دقت بالاتری است.

مراجع

حسینی، ف.، کریمی، ا.، و محمدزاده، م. (۱۳۹۲)، استنباط شبه‌درستنمایی پاسخ‌های فضایی گسسته (مطالعه موردی داده‌های بارندگی استان سمنان)، نشریه علوم دانشگاه خوارزمی، ۳، ۸۰۸-۷۹۷.

Azzalini, A., and Arellano-Valle, R. B. (2013), Maximum Penalized Likelihood Estimation for Skew-Normal and Skew-t Distributions, *Journal of Statistical Planning and Inference*, **143**, 419-433.

- Baghishani, H., and Mohammadzadeh, M. (2011), A Data Cloning Algorithm for Computing Maximum Likelihood Estimates in Spatial Generalized Linear Mixed Models, *Computational Statistics and Data Analysis*, **55**, 1748-1759.
- Bevilacqua, M., Mateu, J., Porcu, E., Zhang, H., and Zini, A. (2010), Weighted Composite Likelihood-Based Tests for Space-Time Separability of Covariance Functions, *Statistics and Computing*, **20**, 283-293.
- Christensen, O. F., and Waagepetersen, R. P. (2002), Bayesian Prediction of Spatial Count Data Using Generalized Linear Mixed Models, *Biometrics*, **58**, 280-286.
- Diggle, P., Tawn, J. A. and Moyeed, R. A. (1998), Model-Based Geostatistic, *Royal Statistical Society, Series C, Applied Statistics*, **47**, 299-350.
- Eidsvik, J., Finley, A. O., Banerjee, S., and Rue, H. (2012), Approximate Bayesian Inference for Large Spatial Datasets Using Predictive Process Models, *Computational Statistics and Data Analysis*, **56**, 1362-1380.
- Firth, D. (1993), Bias Reduction of Maximum Likelihood Estimates, *Biometrika*, **80**, 27-38.
- Green, P. J. (1990). On Use of the EM for Penalized Likelihood Estimation. *Journal of the Royal Statistical Society, Series B(Methodological)*, 443-452.
- Heagerty, P. and Lele, S. (1998), A Composite Likelihood Approach to Binary Spatial Data, *Journal of American Statistical Association*, **93**, 1099-1111.
- Joe, H., and Lee, Y. (2009). On Weighting of Bivariate Margins in Pairwise Likelihood, *Journal of Multivariate Analysis*, **100**, 670-685.
- Lindsay, B. (1988), Composite Likelihood Methods, *Contemporary Mathematics*, **80**, 220-239.
- McCulloch, C. (1997), Maximum Likelihood Algorithms for Generalized Linear Mixed Models, *Journal of the American Statistical Association*, **92**, 162-170.
- McCullagh, P. (1983), Quasi-Likelihood Functions, *Annals of Statistics*, **11**, 59-67.
- McCullagh, C. E., and Nelder, J. A. (1989), *Generalized Linear Models*, Chapman and Hall. London.
- Mohammadzadeh, N. and Hosseini, F. (2011), Maximum-Likelihood Estimation for Spatial GLM Models, *1st Conference on Spatial Statistics 2011*.
- Nelder, J. A., and Wedderburn, R. W. M. (1972), Generalized Linear Models, *Journal of the Royal Statistical Association, Series A*, **135**, 370-384.

- Rue, H., and Martino, S. (2007), Approximate Bayesian Inference for Hierarchical Gaussian Markov Random Field Models, *Journal of Statistical Planning and Inference*, **137**, 3177-3192.
- Tibshirani, R. (1996), Regression Shrinkage and Selection via the Lasso, *Journal of the Royal Statistical Society, Series B*(Methodological), 267-288.
- Torabi, M. (2015), Likelihood Inference for Spatial Generalized Linear Mixed Models, *Communications in Statistics-Simulation and Computation*, **44**, 1692-1701.
- Varin, C., Host, G., and Skare, O. (2005), Pairwise Likelihood Inference in Spatial Generalized Linear Mixed Models, *Computational Statistics and Data Analysis*, **49**, 1173-1191.
- Wedderburn, R. (1974), Quasi-Likelihood Functions, Generalized Linear Models and the Gaussnewton Method, *Biometrika*, **61**, 973-981.



مروری بر نمونه گیری مجموعه رتبه دار در جوامع فضایی

مسعود عظیمیان^۱، محمد مرادی
گروه آمار، دانشگاه رازی کرمانشاه

چکیده: امروزه تحلیل های آمار فضایی به عنوان ابزاری مفید و کارآمد مورد توجه بسیاری از محققین قرار گرفته است. هنگامی که هدف پژوهش بررسی جوامع فضایی است، روش های نمونه گیری کلاسیک از جمله روش نمونه گیری مجموعه رتبه دار، همبستگی فضایی بین واحدها را در نظر نمی گیرند. از این رو انتظار می رود که با تلفیق این روش ها و همبستگی فضایی بین واحدهای نمونه گیری، کارایی این روش ها افزایش یابد. در این مقاله به مروری بر کاربرد نمونه گیری مجموعه رتبه دار در جوامع فضایی پرداخته می شود. مقالاتی که در این زمینه کار شده است و نتایج آن ها ارائه می گردد.

واژه های کلیدی: نمونه گیری مجموعه رتبه دار، نمونه گیری فضایی، همبستگی فضایی، برآوردگر فضایی
کد موضوع بندی ریاضی (۲۰۱۰): 62M30, 62D05

۱ مقدمه

نمونه گیری مجموعه رتبه دار برای اولین بار توسط مک اینتایر (۱۹۵۲) به منظور برآورد میزان محصول علوفه معرفی گردید. این طرح نمونه گیری در مواردی که اندازه گیری متغیر مورد نظر بسیار هزینه بر، زمان بر و گاهی خطرناک است اما رتبه بندی واحدها در مجموعه های کوچکتر آسان است، برآوردگرهای دقیقی در مقایسه با طرح نمونه گیری تصادفی ساده بدون جایگذاری به دست می دهد. در حالت کلی نحوه اجرای این طرح نمونه گیری به این صورت است که ابتدا یک نمونه تصادفی ساده بدون جایگذاری (SRSWOR)^۱ با اندازه m^2 از جامعه ای متناهی با اندازه N واحدی با میانگین $\bar{Y}$ ، و واریانس σ^2 استخراج و بطور تصادفی به m مجموعه هر کدام با اندازه m تقسیم می کنیم. سپس واحدهای هر مجموعه را صورت صعودی بر اساس یک متغیر کمکی یا هر اطلاعات کمکی دیگری مرتب کرده، از مجموعه اول، واحد متناظر با اولین آماره ترتیبی، از مجموعه دوم واحد متناظر با دومین آماره ترتیبی و ... و از مجموعه m ام واحد متناظر با m امین آماره مرتب انتخاب

¹ Simple Random Sampling Without Replacement

^۱ نام و ایمیل ارائه دهنده مقاله: مسعود عظیمیان، azimian.masoud@razi.ac.ir

می‌شود. مراحل با هم را یک چرخه می‌نامند که نتیجه آن یک نمونه مجموعه‌رتبه‌دار به اندازه m است. به منظور دستیابی به یک نمونه با اندازه $n = rm$ ، می‌توان این چرخه را به تعداد r بار تکرار کرد.

هنگامی که با جوامع فضایی کار می‌شود، روش‌های نمونه‌گیری کلاسیک از جمله روش نمونه‌گیری مجموعه‌رتبه‌دار، همبستگی فضایی بین واحدها را در نظر نمی‌گیرند. با در نظر گرفتن جنبه‌ی فضایی (مکانی) واحدها، همبستگی فضایی میان واحدها آشکار می‌شود. بنابراین واحدهایی که در همسایگی هم قرار دارند تمایل به همگنی بیشتری دارند. بنابراین انتخاب یک واحد از این نقاط، انتخاب سایر نقاطی که در یک همسایگی با نقطه انتخاب شده هستند کمک زیادی به شناسایی جامعه هدف نمی‌کند. از این رو انتظار می‌رود با در نظر گرفتن اطلاعات مربوط به همبستگی فضایی واحدهای نمونه‌ای، کارایی روش نمونه‌گیری افزایش یابد. محققانی چون **ساهو و همکاران (۲۰۰۶)**، **اریبا (۱۹۹۳)**، **هدایت و همکاران (۱۹۸۸)** و **کانکیور و ری (۲۰۰۸)** این چنین موقعیت‌هایی را در نظر گرفته و روش‌هایی را برای انتخاب نمونه با استفاده از نمونه‌گیری فضایی ارائه کردند که در آن به واحدهایی که با نمونه‌های قبلی همگنی دارند و در همسایگی آن‌ها هستند احتمال کمی را برای انتخاب شدن نسبت می‌دهد. **اریبا (۱۹۹۳)** روش دنباله‌ای واحد وابسته (DUST) ^۲ را پیشنهاد داد که یک فرآیند انتخاب نمونه برای انتخاب واحدهای فضایی از یک جامعه همبسته‌ی فضایی می‌باشد. در این روش ابتدا یک واحد به تصادف انتخاب می‌شود و باقیمانده واحدهای نمونه‌ای به ترتیب بر اساس وزنی که می‌گیرند، انتخاب می‌شوند به این ترتیب که واحدی که به واحد اولیه نزدیک‌تر است وزن کمتری می‌گیرد، لذا احتمال کمتری را برای انتخاب در نمونه نسبت به واحدهای دورتر که زودتر انتخاب شده‌اند دارد. **اریبا (۱۹۹۳)** به صورت تجربی نشان داد که این تکنیک به طور تقریبی ۳۰٪ از کارایی را نسبت به نمونه‌گیری تصادفی ساده افزایش می‌دهد.

برآوردگرهای فضایی بر اساس همبستگی فضایی واحدهای نمونه‌گیری شده نسبت به برآوردگرهای روش‌های نمونه‌گیری سنتی در جوامع فضایی کاراتر می‌باشند. روش پیش‌بینی (روپال ۱۹۷۰) یکی از تکنیک‌های مفید است که در آن همبستگی فضایی جوامع متناهی برای پیش‌بینی واحدهای مشاهده نشده بر اساس فاصله‌هایشان با واحدهای نمونه‌ای مشاهده شده به اشتراک گذاشته می‌شود. یک راه برای پیش‌بینی واحدهای نامعلوم مشاهده نشده جامعه به وسیله روش وزن‌دهی فاصله‌ای وارون (IDW) ^۳ است که توسط **دونالد (۱۹۶۸)** ارائه شده است.

در این مقاله به مروری بر مقالاتی می‌پردازیم که نمونه‌گیری مجموعه‌رتبه‌دار را در جوامع فضایی مورد بررسی قرار داده‌اند. در بخش دوم مقاله ارائه شده توسط **کانکیور و ری (۲۰۰۸)** با عنوان، نمونه‌گیری مجموعه رتبه‌دار فضایی در جوامع همبسته فضایی، بررسی و نتایج آن ارائه می‌گردد. در بخش سوم مقاله ارائه شده توسط **بیسواز و همکاران (۲۰۱۵)** که به بررسی برآوردگر فضایی براساس نمونه‌گیری مجموعه‌رتبه‌دار در جوامع متناهی دارای همبستگی فضایی پرداخته‌اند ارائه و بررسی می‌شود.

۲ نمونه‌گیری مجموعه‌رتبه‌دار فضایی در جوامع همبسته فضایی

نمونه‌گیری مجموعه‌رتبه‌دار فضایی (SRSS) ^۴ برای اولین بار توسط **کانکیور و ری (۲۰۰۸)** معرفی شد. در واقع در این مقاله یک تکنیک نمونه‌گیری دو مرحله‌ای معرفی می‌شود که در آن نمونه‌گیری مجموعه‌رتبه‌دار با همبستگی فضایی ترکیب می‌شوند. در مرحله اول به منظور تشکیل خوشه‌های فضایی تصادفی (RSC) ^۵، m واحد کلیدی با روش $DUST$ انتخاب می‌شوند. در مرحله دوم مجموعه‌های مرتب‌شده از هر کدام از RSC ها انتخاب می‌گردند. همچنین یک برآوردگر نارایب

^۲Dependent Unit sequential Technique

^۳Inverse Distance Weighting

^۴Spatial Ranked Set Sampling

^۵Random Spatial Clustr

برای این طرح و برای میانگین جامعه براساس برآوردگر راج (۱۹۵۶) ارائه کردند.

۱.۲ نحوه اجرای طرح SRSS

فرض کنید پارامتر موردنظر میانگین متغیر Y جامعه، $\bar{Y}$ و X متغیر کمکی که همبستگی بالایی با متغیر اصلی دارد و برای تمام واحدهای جامعه معلوم است باشند. همچنین فرض کنید که این جامعه دارای N واحد همبسته مکانی باشد. به منظور به دست آوردن یک نمونه به اندازه $n = mr$ مراحل زیر باید اجرا گردد.

۱- ابتدا m واحد به روش $DUST$ انتخاب می شود.

۲- براساس مکان هر کدام از این m واحد، همه واحدهای جامعه به m خوشه فضایی تصادفی (RSC) تقسیم می شوند که هر واحد از جامعه براساس نزدیکی به واحد کلیدی نزدیک به آن، درون یک RSC قرار می گیرد.

۳- درون هر RSC ، r^2 واحد با روش نمونه گیری تصادفی ساده انتخاب کرده و طور تصادفی در r مجموعه دسته بندی می شوند.

۴- واحدهای هر مجموعه براساس صفت کمکی یا براساس قضاوت چشمی به صورت صعودی مرتب می شوند.

۵- پس از مرتب کردن، از مجموعه اول واحد متناظر با اولین اماره ترتیبی، از مجموعه دوم واحد متناظر با دومین اماره ترتیبی ... به همین ترتیب از مجموعه r ام، واحد متناظر با بزرگترین اماره ترتیبی را از نمونه حاصل انتخاب می کنیم.

۶- r واحد انتخاب شده از هر RSC ، یک نمونه مجموعه رتبه دار را تشکیل می دهند.

نمونه حاصل از m خوشه فضایی تصادفی با اندازه $n = r * m$ یک نمونه مجموعه رتبه دار فضایی است.

کانکیور و ری (۲۰۰۸) یک فرمول محاسباتی برای احتمال انتخاب واحدهای نمونه ای را در این طرح نمونه گیری بدست آوردند و براساس آن برآوردگر فضایی مجموعه رتبه دار برای میانگین جامعه را براساس برآوردگر راج (۱۹۵۶) معرفی کردند. نتایج شبیه سازی انجام شده توسط آنها نشان می دهد که طرح نمونه گیری مجموعه رتبه دار فضایی ($SRSS$) ارائه شده که اطلاعات فضایی را در نمونه گیری بکار می گیرد کاراتر از نمونه گیری مجموعه رتبه دار معمولی و نمونه گیری تصادفی ساده است.

۳ روش برآورد فضایی میانگین جامعه تحت طرح نمونه گیری مجموعه رتبه دار

در این بخش مقاله ارائه شده توسط **بیسواز و همکاران (۲۰۱۵)** با عنوان روش برآورد فضایی میانگین جامعه تحت طرح نمونه گیری مجموعه رتبه دار معرفی می گردد. در این مقاله طرح نمونه گیری مجموعه رتبه دار تصادفی شده را برای انتخاب نمونه از یک جامعه فضایی مورد استفاده قرار داده و برآوردگر فضایی مطابق با آن را برای برآورد میانگین جامعه با استفاده از روش IDW از طریق رویکرد پیش بینی معرفی می گردد.

۱.۳ برآوردگر فضایی در حالت نمونه گیری مجموعه رتبه دار

پاتیل و همکاران (۱۹۹۵) نشان دادند که برآوردگر طرح نمونه گیری مجموعه رتبه دار میانگین جامعه، $\bar{Y}$ ، برابر است با:

$$\bar{Y} = \frac{1}{mr} \sum_{k=1}^r \sum_{i=1}^m y_{(i:m)k} \quad (۱.۳)$$

که در آن $y_{(i:m)k}$ مقدار واحد مرتب شده i ام در مجموعه k ام از چرخه k ام است. فرض کنید که $d_{(i:m)k,j}$ فاصله واحد $(i:m)k, j$ از نمونه و هر واحد غیر نمونه ای U_j ، $j \in \bar{s}$ ، باشد که در آن $\bar{s}$ مجموعه همه واحدهای غیر نمونه ای جامعه است که در نمونه انتخاب نشده اند. باتوجه به اینکه واحد نمونه ای $(i:m)k, j$ به صورت تصادفی انتخاب می شود،

$d_{(i:m)k,j}$ مقداری تصادفی است. بنابراین مقدار پیش‌بینی شده‌ی Z_j از میان واحد غیر نمونه‌ای، $j \in \bar{s}$ ، با پیروی از دونالد (۱۹۶۸) از رابطه زیر بدست می‌آید.

$$y_{j,p} = \frac{\sum_{k=1}^r \sum_{i=1}^m \frac{y_{(i:m)k}}{d_{(i:m)k,j}}}{\sum_{k=1}^r \sum_{i=1}^m \frac{1}{d_{(i:m)k,j}}}, \quad (i:m)k \in s, j \in \bar{s} \quad (2.3)$$

با ترکیب کردن این مقادیر پیش‌بینی شده $N - mr$ واحد غیر نمونه‌ای میانگین این واحدها به صورت زیر بدست می‌آید:

$$\bar{y}_{p,RSS} = \frac{1}{N - mr} \sum_{j \in \bar{s}} y_{j,p} \quad (3.3)$$

بنابراین با توجه به روش پیش‌بینی رویال (۱۹۷۰) برآوردگر فضایی تحت طرح نمونه‌گیری مجموعه‌رتبه‌دار ارائه شده توسط بیسواز و همکاران (۲۰۱۵) برای پارامتر میانگین جامعه به صورت زیر می‌باشد:

$$\hat{Y}_{SE,RSS} = \frac{mr\bar{y}_{RSS} + (N - mr)\bar{y}_{p,RSS}}{N} \quad (4.3)$$

نتایج شبیه‌سازی شده انجام شده توسط آن‌ها نشان داد که برآوردگر فضایی ارائه شده بر اساس نمونه‌گیری مجموعه‌رتبه‌دار با توجه به اینکه همبستگی فضایی واحدهای جامعه را در برآورد میانگین جامعه در نظر می‌گیرد کاراتر از برآوردگر معمول نمونه‌گیری مجموعه‌رتبه‌دار و نمونه‌گیری تصادفی ساده است.

مراجع

- Arbia, G. (1993), The Use of GIS in Spatial Statistical Surveys, *International Statistical Review*, **61**, 339-359.
- Biswas, A. , Rai, A. Ahmad, T. and Sahoo, P. M. (2015), Spatial Estimation Approach Under Ranked Sst Samling from Spatial Correlation Finite Population, *Journal of the Indian Society of Agricultural Statistics*, **11**, 551-558.
- Donald, S. (1968), A Two-Dimentional Interpolation Function for Irregularly-Spaced Data, *Proceedings of the 1968 ACM National Conference*, 517-524.
- Hedayat, A. S., Rao, C. R. and Stufken, J. (1988), Sampling Plans Excluding Contiguous Units, *Journal of Statistical Planning and Inference*, **19**, 159-170.
- Kankure, A. K. and Rai, A. (2008), Spatial Ranked Set Sampling from Spatially Correlated Population, *Journal of the Indian Society of Agricultural Statistics*, **62**, 221-230.
- McIntyre, G. A. (1952), A Mmthod of Unbiased Selective Samling Using Ranked Sets, *Australian Journal of Agricultural Research*, **3**, 387-390.
- Patil, G. P., A. K. Sinha and C. Taillie (1995), Finite Population Corrections for Ranked Sset Sampling, *Annals of the Institute of Statistical Mathematics*, **47**, 621-636.

Royall, R. M. (1970), On Finite Population Sampling Theory Under Certain Linear Regression Models, *Biometrika*, **57**, 377-387.

Raj, D. (1956), Some Estimators With Varying Probabilities Without Replacement, *Journal of the American Statistical Association*, **51**, 269-284.

Sahoo, P. M., Singh, R. and Rai A. (2006), Spatial Sampling Procedure for Agricultural Surveys Using Geographical Information System, *Journal of the Indian Society of Agricultural Statistics*, **60(2)**, 134-143.



پیشگویی فضایی سه بعدی به منظور برآورد مقاومت فشاری تک محوری در یکی از مخازن نفتی ایران

فتانه فخری^۱، حسین باغیشنی^۱، بهزاد تخمچی^۲
^۱گروه آمار، دانشکده علوم ریاضی، دانشگاه صنعتی شاهرود
^۲گروه اکتشاف معدن، دانشکده مهندسی نفت، معدن و ژئوفیزیک

چکیده: روش‌های مدل‌سازی مبتنی بر زمین‌آمار نقش مهمی را در صنعت اکتشاف معدن و نفت ایفا می‌کند. در این صنایع با در نظر گرفتن هزینه‌های زیادی که صرف حفاری می‌شود، شناخت ناحیه مورد مطالعه و پیش‌گویی فضایی سه‌بعدی آن می‌تواند به کم کردن هزینه‌ها و کاهش خطا کمک کند. در این مقاله تغییرنگار سه بعدی و پیش‌گویی فضایی را معرفی کردیم و برای تشریح روش از مجموعه داده ژئومکانیکی مقاومت فشاری تک محوری استفاده کردیم.
واژه‌های کلیدی: پیش‌گویی فضایی سه بعدی، زمین‌آمار، مقاومت فشاری تک محوری
کد موضوع‌بندی ریاضی (۲۰۱۰): 62H11, 97K80.

۱ مقدمه

دنیایی که ما در آن زندگی می‌کنیم مملو از وابستگی‌هاست، در تحقیقات نادیده گرفتن این وابستگی می‌تواند نتایج غلط و زیان‌های را بر پژوهشگر و کاربران آن عرصه متحمل سازد. در علوم مختلف از جمله نفت، زمین‌شناسی، بوم‌شناسی گیاهی و حیوانی، سنجش از دور^۱، پزشکی و ... داده‌هایی که استفاده می‌شود به موقعیت مکانی و جهت قرارگیری از هم بستگی دارد، بطوری که در اغلب موارد هر چه فاصله بین مشاهدات کمتر باشد شباهت بین آن‌ها بیشتر و در فاصله بیشتر شباهت کمتر است. تجزیه و تحلیل این نوع داده‌ها که به داده‌های فضایی شهرت دارند با روش‌های استنباطی آمار کلاسیک ناممکن است. وارد کردن ساختار وابستگی القا شده توسط موقعیت‌های مکانی داده‌ها (ساختار وابستگی فضایی) در تحلیل داده‌ها توسط روش‌های آمار فضایی ممکن شده است. از آمار فضایی در بسیاری از روش‌ها و مدل‌بندی‌های آماری برای داده‌های فضایی استفاده می‌شود (کرسبی، ۱۹۹۳) یکی از زیرشاخه‌های آمار فضایی زمین‌آمار است و به مدل‌ها و روش‌هایی اشاره دارد که از

^۱ Remote Sensing

دو مشخصه پیروی کند. اول اینکه Y_i ها $i = 1, \dots, n$ مشاهداتی هستند از مجموعه گسسته مشتمل بر مکان نمونه‌های Z_i که در برخی از مناطق فضایی مانند A مشاهده می‌شوند، و دوم هر مقدار مشاهده شده Y_i با اندازه‌گیری مستقیم و یا با روابط آماری برآورد شده باشد (دیگل و ریبریو، ۲۰۰۷). داده‌های زمین‌آماری در موقعیت‌های ثابت و در ناحیه‌ای پیوسته مشاهده می‌شوند و متغیر مورد بررسی می‌تواند پیوسته و یا گسسته باشد. در حقیقت زمین‌آمار از ابتدای دهه ۱۹۵۰ به بعد در صنعت معدن توسعه یافت که در حال حاضر بطور گسترده‌ای در سایر علوم زیست محیطی برای نقشه‌برداری، نظارت و مدیریت استفاده می‌شود. همچنین با استفاده از مدل‌سازی همبستگی فضایی، زمین‌آمار می‌تواند پیش‌بینی‌های نارایب با کمترین واریانس را ارائه دهد (اولیور و وستر، ۲۰۱۵). روش‌های مدل‌سازی مبتنی بر زمین‌آمار نقش مهمی را در صنعت اکتشاف (معدن و نفت) ایفا می‌کند. چرا که در این صنایع با در نظر گرفتن هزینه‌های زیادی که صرف حفاری می‌شود، شناخت ناحیه مورد مطالعه و پیشگویی سه‌بعدی آن می‌تواند به کم کردن هزینه‌ها و کاهش خطا کمک کند. پیش‌بینی ویژگی‌های الاستیک ژئومکانیکی مخازن نفتی، به‌طور گسترده، نقش مهمی در بهینه‌سازی و کاهش خطرات و افزایش بهره‌وری حفاری در یک مخزن ایفا می‌کند (نظری استاد و همکاران، ۲۰۱۸). هر چاه دارای طول و عرض جغرافیایی است و از طرفی هزینه‌های حفاری با افزایش عمق به صورت نمایی رشد می‌کند (عبدالآل، ۲۰۱۳). مدل‌سازی سه‌بعدی متغیرهای تاثیرگذار در حفاری می‌تواند علاوه بر افزایش سرعت و کاهش هزینه، منجر به برنامه‌ریزی بهتر فرآیند حفاری و کاهش مخاطرات طبیعی گردد. با این حال، ساخت مدل‌های سه بعدی قابل اعتماد برای این ویژگی‌ها همچنان با سوال‌ها و گمانه‌هایی همراه است که پاسخ‌گویی به آن‌ها مستلزم تحقیق و توسعه مدل‌های کارا در سه بعد است.

در این پژوهش در بخش ۳ و ۴ به بررسی ناحیه مورد مطالعه می‌پردازیم و نیز تعاریف تغییرنگار و پیش‌گوی فضایی را بازگو می‌کنیم، سپس به پیشگویی سه بعدی مقاومت فشاری تک محوری می‌پردازیم و در نهایت در بخش پایانی به بحث و نتیجه‌گیری خواهیم پرداخت.

۲ ناحیه تحت مطالعه

داده‌های استفاده شده در این پژوهش مربوط به مقاومت فشاری تک محوری سنگ^۲ (UCS) است. یکی از پارامترهای مهم در ارزیابی مقاومت سنگ و ساخت مدل ژئومکانیک، مقاومت تک محوری سنگر بکر UCS است. معمولاً مقدار مقاومت فشاری از طریق آزمایش بدست می‌آید از طرفی دسترسی به مغزه حفاری مشکل است لذا محققان صنعت نفت روابطی را جهت برآورد پارامترهای مکانیکی از روی نگاره‌های پتروفیزیکی ارائه داده‌اند. در این پژوهش برای محاسبه مقدار مقاومت فشاری تک محوری از این روابط استفاده شده است. منطقه مورد مطالعه یکی از مخازن نفتی ایران شامل ۱۰ حلقه چاه با طول و عرض جغرافیایی و عمق مشخص است که مقدار UCS با استفاده از متغیرهای ثبت شده مرتبط با این متغیر بطور تقریبی محاسبه شده است. لازم به ذکر است که چاه‌ها دارای ضخامت‌های متفاوتی هستند، از طرفی این مخزن بر روی یک چین زمین‌شناسی قرار دارد پس از رابطه (۱.۲) برای یکسان‌سازی عمق چاه‌ها در همان مختصات طول و عرض جغرافیایی استفاده شد.

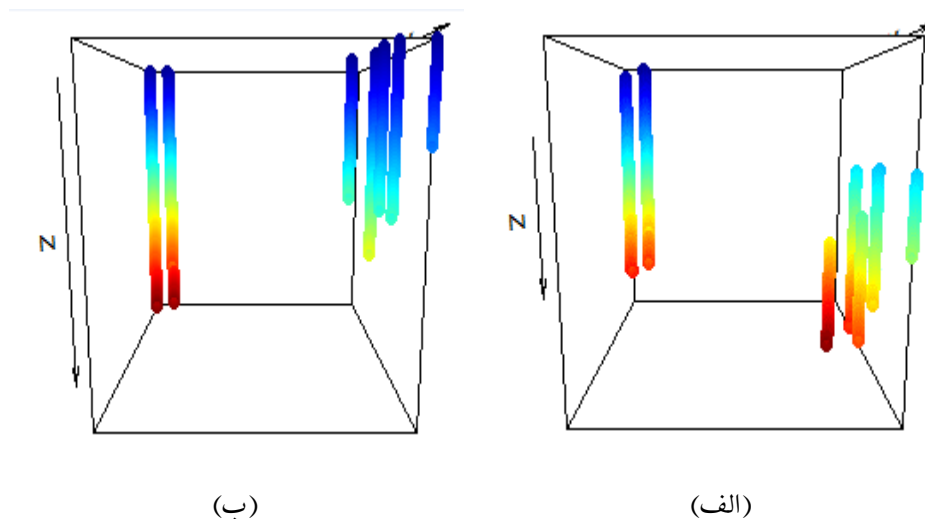
$$z_{(new)i} = z_{(old)i} - |\min(z_{ij}) - \min(z_{(old)i})|, \quad j = 1 \dots 10, \quad i = 1, \dots, m_j \quad (1.2)$$

که $z_{(old)i}$ ، $z_{(new)i}$ به ترتیب عمق حقیقی و اصلاح شده برای مکان i ام در هر چاه است، $\min(z_{ij})$ کمترین عمق اندازه‌گیری شده در تمام چاه‌ها و m_j نشان دهنده تعداد مشاهدات در هر چاه است. در شکل ۱ نمودار سه بعدی چاه‌ها با

²Unconfined Compressive Strength

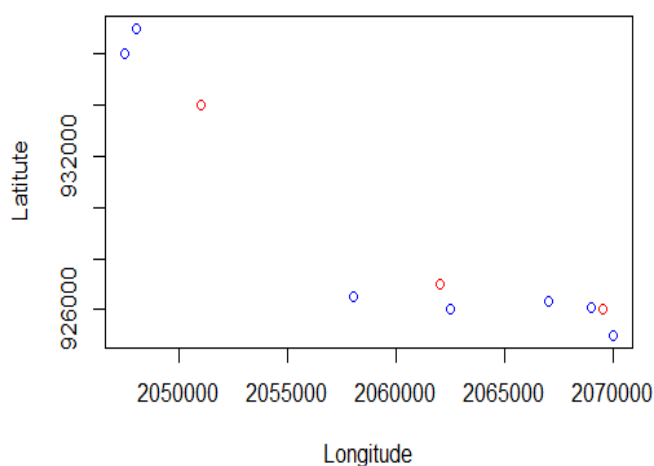
³Test set

⁴Cross validation



شکل ۱: توضیحات نمودار الف: با عمق حقیقی ب: بعد از یکسان سازی عمق

عمق حقیقی و عمق اصلاحی آورده شده است. از بین چاه‌ها سه چاه که شامل ۸۸۷۳ مشاهده است را بطور تصادفی برای مجموعه داده‌های آزمایشی^۳ انتخاب کردیم تا نتایج را بر اساس اعتبار سنجی متقابل^۴ مورد ارزیابی قرار دهیم و از بقیه چاه‌ها که ۲۸۴۵۵ مشاهده را شامل می‌شود بعنوان مجموعه آموزشی^۵ برای مدل‌سازی سه بعدی استفاده کردیم. شکل ۲ مربوط به محل قرارگیری چاه‌ها در موقعیت‌های مکانی طول و عرض جغرافیایی به تفکیک مجموعه آموزشی و آزمایش است.



شکل ۲: محل قرارگیری چاه‌های مجموعه آموزشی (سیاه) و مجموعه آزمایشی (قرمز) در دو بعد

⁵Train set

۳ تغییرنگار

در آمار فضایی تغییرنگار نظری $\gamma(h)$ واریانس تفاضل بین دو مقدار مشاهده شده میدان تصادفی در دو مکان فضایی که به فاصله h از هم قرار دارند و بصورت

$$\gamma(h) = \text{Var}(Z(s+h) - Z(s)) \quad (۱.۳)$$

$\gamma(h)$ تعریف می‌شود، نیم‌تغییرنگار^۶ نامیده می‌شود. تابعی برای توصیف میزان وابستگی مقادیر میدان تصادفی و یا فرآیند برآورد دقیق تغییرنگار برای پیش‌گویی قابل اعتماد توسط کریگینگ مورد نیاز است. تغییرنگار تجربی (نمونه) معمولاً توسط روش‌های گشتاوری در دنباله‌ای از گام‌ها محاسبه می‌شود و یک یا چند مدل معتبر (مجاز^۷) تغییرنگار بر روی آن برازش داده می‌شود. یک تغییرنگار ممکن است در طول یک برش یا توری در فواصل منظم یا در فواصل پراکنده نامنظم محاسبه شود. دقت تغییرنگار به اندازه نمونه، تعداد گام‌هایی که برآورد شده است و نیز به فاصله بازه نسبت به مقیاس فضایی یا متغیر، توزیع کناری، ناهمسانگردی و روند بستگی دارد. برآوردگرهای تنومند^۸ نسبت به مقادیر دورافتاده^۹ حساس نیستند. تغییرنگار ممکن است دارای محدودیت (کراندار) باشد (برای مانای مرتبه دوم^{۱۰}) و یا بدون کران (تنها در مانای ذاتی^{۱۱}). پارامترهای تغییرنگار که توسط مدل‌های متغیر تغییرنگار برآورد شده‌اند خلاصه‌ای از تغییرات فضایی را که برای کریگینگ نیاز است را به ما می‌دهند. در صورت وجود ناهمسانگردی در میدان تصادفی محاسبه تغییرنگار می‌تواند حداقل در سه جهت ناهمسانگردی را تعیین کند. میانگین حداقل مربعات باقی‌مانده‌ها^{۱۲} و معیار آکایک^{۱۳} می‌تواند در انتخاب بهترین مدل برازش داده شده به ما کمک کند. (اولیور و وستر، ۲۰۱۵) برآورد تغییرنگار معمولاً براساس روش گشتاوری ماترون^{۱۴} (MOM) به صورت

$$\hat{\gamma}(\mathbf{h}) = \frac{1}{m(\mathbf{h})} \sum_{i=1}^{m(\mathbf{h})} \{z(\mathbf{x}_i) - z(\mathbf{x}_i + \mathbf{h})\}^2 \quad (۲.۳)$$

محاسبه می‌گردد. که $z(\mathbf{x}_i)$ و $z(\mathbf{x}_i + \mathbf{h})$ مقادیر مشاهده شده z در موقعیت‌های $\mathbf{x}_i$ و $\mathbf{x}_i + \mathbf{h}$ و $m(\mathbf{h})$ تعداد جفت نمونه‌هایی است که در گام $\mathbf{h}$ قرار دارد. محاسبه تغییرنگار سه‌بعدی در میدان با توری‌های منظم یا نامنظم به این شکل است که می‌توان توری را بصورت یک سری در نظر گرفت و در سه بعد مورد بررسی قرار داد. معادله (۳.۳) مربوط به تغییرنگار سه بعدی است که برای بررسی همسانگردی باید در چند جهت مختلف در طول ردیف، ستون و روی قطر یک توری محاسبه شود.

$$\hat{\lambda}(p, q) = \frac{1}{m(p, q)(n - q)(l - r)} \sum_{i=p}^{m-p} \sum_{j=q}^{n-q} \sum_{k=r}^{l-r} \{z(i, j, k) - z(i + p, j + q, k + r)\}^2 \quad (۳.۳)$$

که p, q, r به ترتیب اندازه گام و m, n, l تعداد گام‌های ردیف، ستون و بعد سوم (در اینجا عمق) هستند. نمودار سه‌بعدی نیم‌تغییرنگار تجربی داده‌های مقاومت فشاری در چهار جهت ۰، ۴۵، ۹۰ و ۱۳۵ درجه برای بررسی ناهمسانگردی

^۶Semivariogram

^۷Authorized

^۸Robust estimators

^۹Outliers

^{۱۰}Second-order stationary

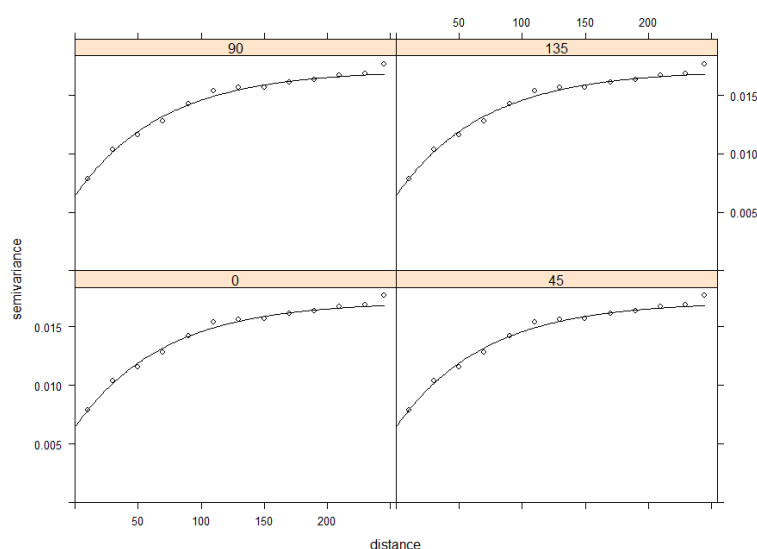
^{۱۱}Intrinsic stationary

^{۱۲}Residual mean squares

^{۱۳}Akaike

^{۱۴}Matheron's method of moments

همراه با منحنی مدل پارامتری نمایی برازش داده شده در شکل ۳ آورده شده است. با توجه به شکل می‌توان دریافت که نیم‌تغییرنگار در این چهار جهت همسان‌گرد است.



شکل ۳: نمودار سه‌بعدی نیم‌تغییرنگار و منحنی مدل پارامتری نمایی برازش داده شده در چهار جهت

۴ پیش‌گویی فضایی

یکی از اهداف اصلی در تحلیل داده‌های فضایی، پیش‌گویی پاسخ در مکان‌های فاقد مشاهده است که با استفاده از روش کریگینگ انجام می‌شود (کرسبی، ۱۹۹۳). اصطلاح کریگینگ به افتخار دانیال کریگ^{۱۵} (۲۰۱۳ - ۱۹۱۹) نام‌گذاری شده است، کریگ در استفاده از روش‌های آماری برای ارزیابی و تحلیل مخازن معدنی پیشگام بود و بعدها این روش‌ها توسط ماترون^{۱۶} و دیگر همکارانش در دانشکده معدن پاریس فرموله، توسعه و نشر داده شد. (دیگل و جورجی، ۲۰۱۹). کریگینگ بهترین پیش‌گویی خطی نقطه‌ای و بلوکی است، به‌طوری که بهترین در این‌جا به معنی پیش‌گو با کمترین واریانس است. پیش‌گویی کریگینگ یک میانگین وزنی است که به نقاط نزدیک‌تر وزن بیشتر و به نقاط دورتر وزن کمتری اختصاص می‌دهد (اولیور و وستر، ۲۰۱۵). یکی از پذیره‌های روش کریگینگ، مانایی میدان تصادفی است. با معلوم بودن ساختار همبستگی فضایی و در نظر گرفتن تابع زیان توان دوم خطا، پیش‌گویی بهینه (کریگینگ) از می‌نیم کردن میانگین توان دوم خطای پیش‌گو به دست می‌آید. در روش‌های کلاسیک، میدان تصادفی را گاوسی در نظر می‌گیرند. بر حسب ساختار فضایی میدان تصادفی، سه نوع کریگینگ تعریف می‌شوند: (۱) کریگینگ ساده^{۱۷} که در آن فرض می‌شود میانگین میدان معلوم است؛ (۲) کریگینگ عادی^{۱۸} که در آن فرض می‌شود میانگین ثابت ولی نامعلوم است؛ و (۳) کریگینگ عام^{۱۹} که در آن میانگین به عنوان تابعی از موقعیت‌ها (و سایر متغیرهای تبیینی) در نظر گرفته می‌شود و به آن روند می‌گویند. در ادامه، به‌طور مختصر، کریگینگ ساده و عادی را معرفی می‌کنیم.

¹⁵Daniel G. Krige

¹⁶Matheron

¹⁷Simple Kriging

¹⁸Ordinary Kriging

¹⁹Universal Kriging

۱.۴ کریگینگ عادی

فرض کنید مقادیر متغیر پاسخ Z را در n مکان ناحیه تحت مطالعه، $D \subset \mathbb{R}^d$ ، مثل $\{s_1, \dots, s_n\}$ ، مشاهده کرده باشیم و هدف پیش‌گویی مقدار پاسخ در مکان جدید s_0 ، یعنی $Z(s_0)$ ، باشد. در کریگینگ عادی تغییرات مشاهدات پاسخ از نظر فضایی به هم وابسته و تصادفی و تنها در یک مقدار نامعلوم μ ، مشترک هستند. در این حالت می‌توان میدان تصادفی گاوسی $\{Z(s) : s \in D\}$ را به صورت زیر تجزیه کرد:

$$Z(s) = \mu + \delta(s)$$

که در آن μ ثابت و نامعلوم و $\delta(\cdot)$ میدان تصادفی مانای گاوسی با میانگین صفر است. کریگینگ عادی، به عنوان یک میانگین وزنی از مشاهدات، به صورت

$$\hat{Z}(s_0) = \sum_{i=1}^n \lambda_i Z(s_i)$$

تعریف می‌شود. ضرایب λ_i وزن‌های کریگینگ هستند که سهم هر کدام از مشاهدات را (بسته به فاصله آن‌ها تا مکان جدید s_0) در پیش‌گویی پاسخ یدک می‌کشند. چنانچه شرط $\sum_{i=1}^n \lambda_i = 1$ را در نظر بگیریم، کریگینگ یک پیش‌گوی ناریب خواهد بود، یعنی $E[\hat{Z}(s_0)] = \mu$. مقادیر بهینه ضرایب $\lambda = (\lambda_1, \lambda_2, \dots, \lambda_n)'$ از حل معادله

$$\lambda' = (\gamma + \mathbf{1}_n' \frac{(\mathbf{1} - \mathbf{1}_n' \mathbf{\Gamma}^{-1} \gamma)}{\mathbf{1}_n' \mathbf{\Gamma}^{-1} \mathbf{1}_n})' \mathbf{\Gamma}^{-1} \quad (1.4)$$

به دست می‌آیند، که در آن $\mathbf{\Gamma}$ یک ماتریس $n \times n$ با درایه (i, j)

$$\gamma(s_i - s_j) = \frac{1}{\gamma} \text{Var}(Z(s_i) - Z(s_j))$$

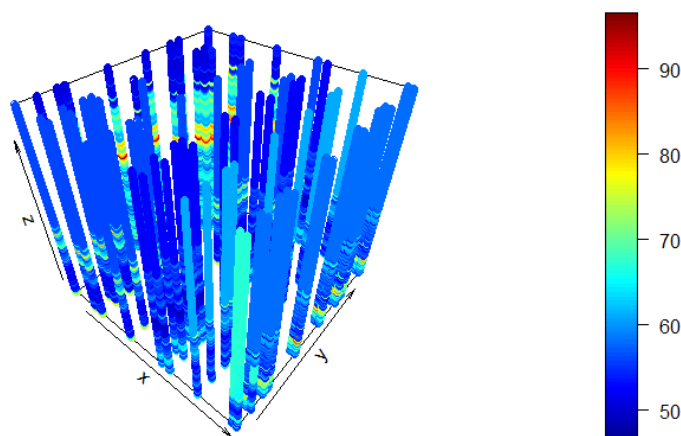
و $\gamma = (\gamma(s_0 - s_1), \gamma(s_0 - s_2), \dots, \gamma(s_0 - s_n))'$ ، به طوری که $\gamma(\cdot)$ تابع نیم‌تغییرنگار است. همچنین $\mathbf{1}_n$ برداری به طول n از ۱‌ها است. در مورد چگونگی رسیدن به این ضرایب و نحوه محاسبه واریانس کریگینگ عادی می‌توانید به **کرسی، (۱۹۹۳)** مراجعه کنید.

کریگینگ ساده حالت خاصی از کریگینگ عادی است که در آن فرض می‌شود، بین مشاهدات روندی وجود ندارد و میانگین معلوم و مستقل از مختصات فضایی است. در این حالت شرط ناریبی صفر بودن مجموع ضرایب لازم نیست و معادله محاسبه ضرایب کریگینگ، مشابه (۱.۴)، با کمی تغییرات جزئی نتیجه می‌شود.

۵ پیش‌گویی سه بعدی مقاومت فشاری تک محوری سازند

برای پیش‌گویی داده‌های مقاومت فشاری در سایر نقاط ناحیه تحت مطالعه در شکل ۴ با توجه به پذیره گاوسی بودن میدان، با استفاده از کریگینگ عادی و شبکه‌بندی ناحیه تحت مطالعه، نقشه پهنه‌بندی پیش‌گویی شده مقدار مقاومت فشاری و واریانس پیش‌گویی را به دست آوردیم. شکل ۵ نمودار دو بعدی مقادیر پیش‌گویی و واریانس پیش‌گویی را در سه عمق نشان می‌دهند.

محور افقی شکل ۶ نشان دهنده مقادیر واقعی مجموعه داده‌های آزمایشی Z_i و محور عمودی مقادیر پیش‌گویی بدست آمده از کریگینگ عادی $\hat{Z}_i$ برای مجموعه آزمایشی است. برای UCS کمتر از ۸۰ توزیع پراکنش داده‌ها حول نیم‌ساز ربع



شکل ۴: پیشگویی سه بعدی میدان نفتی تحت مطالعه

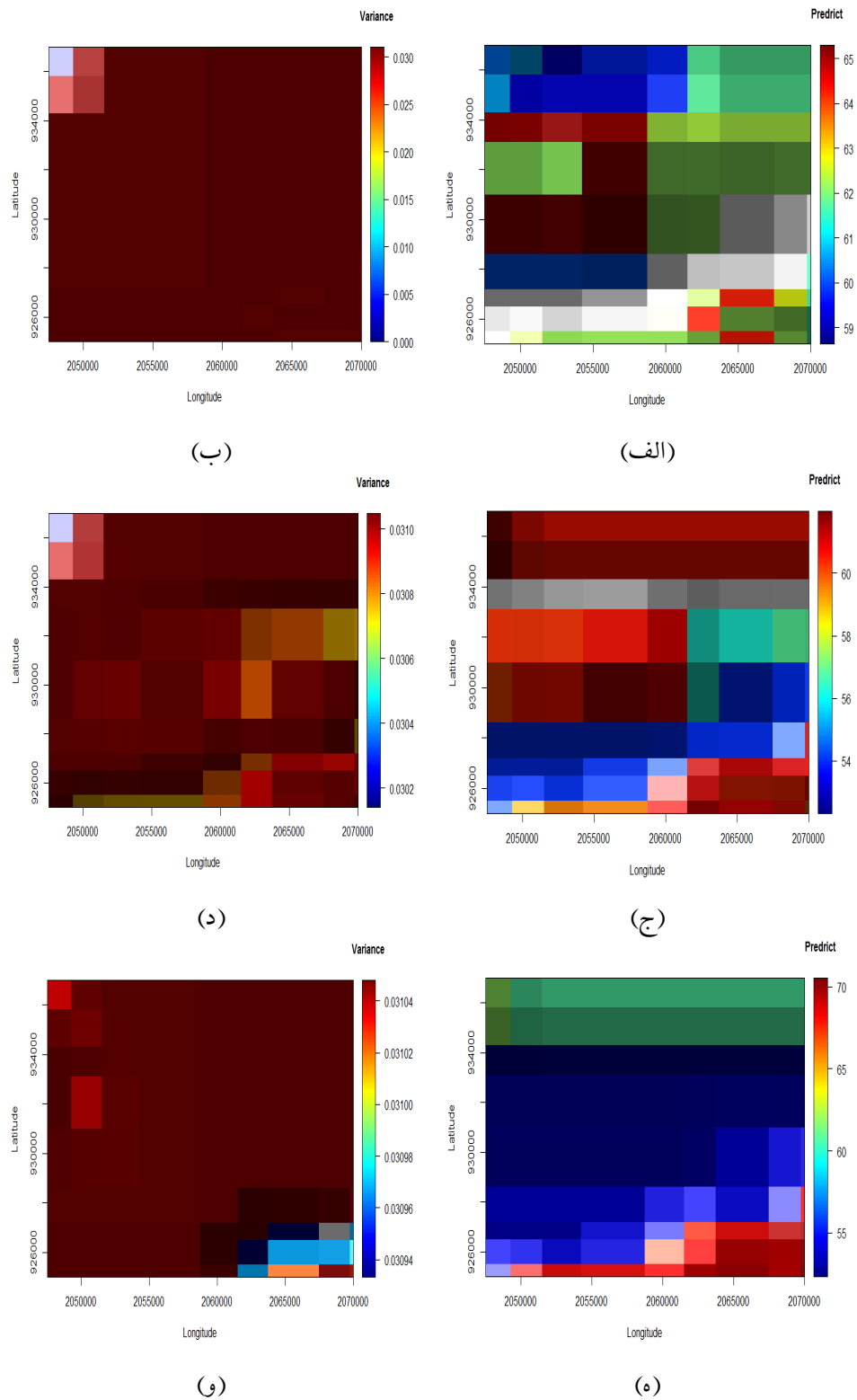
اول و سوم ($Z_i = \hat{Z}_i$) است و بیانگر مناسب بودن مدل و مقدار ناچیز باقی مانده‌ها است. با افزایش مقادیر داده‌ها، پیشگو و مقادیر اصلی از نیم‌ساز فاصله گرفته و اندازه خطا بیشتر شده است، به عبارتی واریانس خطا ثابت نیست ولی همچنان نیم‌ساز بین داده‌ها قرار دارد و میانگین خطا برابر با صفر است.

بحث و نتیجه‌گیری

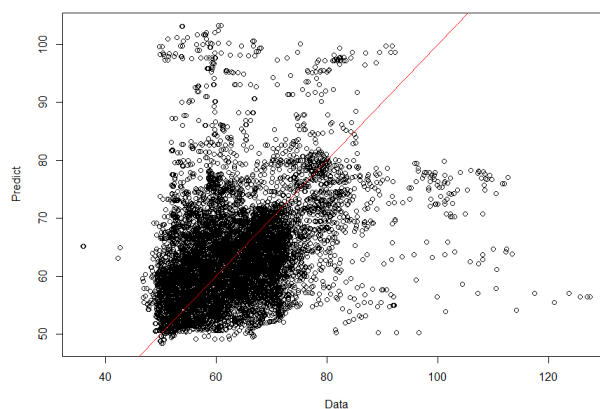
در بسیاری از علوم از داده‌های سه‌بعدی برای تحلیل و مدل‌سازی استفاده می‌شود و بررسی مدل‌سازی و پیشگویی در سه بعد می‌تواند به روشن شدن بعضی از ابهامات کمک کند. در این پژوهش از کریگینگ عادی سه بعدی برای پیشگویی فضایی مقادیر مقاومت فشاری تک محوری استفاده شد. ابتدا به طور تصادفی ۷۰ درصد از کل داده‌ها به عنوان مجموعه آموزشی در نظر گرفته شد و با برازش مدل مناسب به نمودار نیم‌تغییرنگار سه بعدی داده‌ها پیشگویی ناحیه تحت مطالعه صورت گرفت و در آخر نمودار پراکنش داده‌های مجموعه آزمایشی در مقابل مقادیر پیشگویی رسم شد که می‌توان گفت مناسب بودن مدل و مقدار ناچیز باقی مانده‌ها را نشان می‌دهد.

مراجع

- محمدزاده، م.، (۱۳۹۸)، آمار فضایی و کاربردهای آن، چاپ سوم، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران.
- Abdel-Aal, H. K., Alsahlawi, M. A. (Eds.), (2013), *Petroleum Economics and Engineering*, CRC Press.
- Cressie, N. (1993), *Statistics for Spatial Data*, Hoboken, NJ, USA: John Wiley and Sons, Inc.
- Diggle, P. J., Ribeiro, P. J. (2007), *Model Based Geostatistics*, Springer Series in Statistics.
- Diggle, P. J., Giorgi, E. (2019), *Model Based Geostatistics for Global Public Health: Methods and Applications*, Chapman and Hall/CRC.



شکل ۵: نقشه پهنه بندی پیش‌گویی مقادیر مقاومت فشاری در عمق‌های ۱۷۰۳، ۲۱۵۵ و ۲۴۸۰ متر به ترتیب در (الف)، (ج) و (هـ). همچنین نقشه پهنه‌بندی واریانس پیش‌گویی در همین عمق‌ها به ترتیب در (ب)، (د) و (و).



شکل ۶: مقادیر پیشگویی در مقابل مقادیر واقعی داده‌های آزمایشی

Oliver, M. A., Webster, R. (2015), *Basic Steps in Geostatistics: the Variogram and Kriging* (Vol. 106), New York: Springer

Ostad, M. N., Niri, M. E., Darjani, M. (2018), 3D Modeling of Geomechanical Elastic Properties in a Carbonate-Sandstone Reservoir: A Comparative Study of Geostatistical Co-Simulation Methods, *Journal of Geophysics and Engineering*, **15**, DOI: 10.1088/1742-2140/aaa983.



بررسی اثر موقعیت مکانی هتل‌ها بر میزان رضایت‌مندی مشتریان

زهرا فضلعلی^۱، کیومرث مترجم

گروه آمار، دانشگاه تربیت مدرس

چکیده: در این مقاله اثر موقعیت مکانی هتل‌ها بر میزان رضایت‌مندی مشتریان در ایالت کالیفرنیا، آمریکا طی سال‌های ۲۰۰۲ تا ۲۰۱۵ بررسی شده است. هدف از انجام این مطالعه، تحقق هر چه بیشتر و بهتر خواسته‌های مشتریان و همچنین دستیابی به موقعیت مکانی مناسب جهت ایجاد هتل بوده است. جامعه آماری این تحقیق شامل همه نظرات ثبت شده‌ای است که در مورد هتل‌های ایالات متحده در برنامه TripAdvisor.com توسط کاربران ثبت و جمع‌آوری شده است. با در نظر گرفتن اطلاعات موقعیت جغرافیایی، این مطالعه به بررسی ۲۲۰ پاسخ درباره هتل‌ها از مجموع ۷۷۰۹۸ نظری که توسط کاربران ثبت شده پرداخته است. با استفاده از تحلیل اکتشافی نشان خواهیم داد که بین رضایت مشتری و موقعیت مکانی هتل‌ها همبستگی فضایی وجود دارد. نتایج این تحقیق در حوزه‌های بازاریابی و رویکردهای مدیریتی صنعت هتل‌داری بسیار کارآمد و مورد استفاده مدیران این حوزه می‌باشد. **واژه‌های کلیدی:** آمار فضایی، رضایت‌مندی مشتری، بازاریابی. کد موضوع‌بندی ریاضی (۲۰۱۰): 91D25، 62M30، 62H11.

۱ مقدمه

برای نشان دادن همبستگی فضایی بین رضایت‌مندی مشتریان از هتل‌ها در سطح شهر سانفرانسیسکو از آزمون موران استفاده می‌کنیم. از آنجا که رضایت مشتری در رویکردهای مدل‌سازی تجربی معمولاً به عنوان متغیر وابسته رفتار می‌کند لذا وجود همبستگی فضایی در این متغیر باعث می‌شود استفاده از مدل‌های معمول در آمار، نتایج مطلوبی به همراه نداشته باشد. بنابراین، برای تجزیه و تحلیل اینگونه داده‌ها به یک الگوی فضایی نیاز پیدا خواهیم کرد. در بخش ۲ این مقاله، داده‌ها و رویکردهای پاکسازی آنها معرفی می‌شود. در بخش ۳ همبستگی فضایی رضایت‌مندی مشتریان، آماره آزمون موران و ماتریس‌های همسایگی معرفی می‌شوند و در بخش آخر به بحث و نتیجه‌گیری می‌پردازیم. رضایت مشتری فاکتوری مهم در موفقیت سازمان‌ها در محیط تجاری است. همه سازمان‌ها می‌کوشند رضایت مشتری را از طریق بهبود کیفیت خدمات یا محصولاتشان افزایش دهند. اما در بسیاری از مواقع این سازمان‌ها به علت کمبود منابع و تصمیم‌گیری نادرست درباره

^۱ نام و ایمیل ارائه دهنده مقاله: زهرا فضلعلی، Zahrafazlali467@gmail.com

اینکه چه مجموعه‌ای از فعالیت‌ها سبب بهترین عملکرد در بهبود رضایت مشتری می‌شود معمولاً به اهداف مشخص شده خود دست پیدا نمی‌کنند. این امر به این دلیل رخ می‌دهد که اساساً بعضی از عوامل نقش اصلی را در رضایت مشتریان یک سازمان ایفا نمی‌کنند، لذا شناسایی عوامل حیاتی و تصمیم‌ساز در رضایت مشتریان در موفقیت هر سازمان یا شرکت، نقشی اساسی و ضروری دارد. یکی از مهمترین تحولاتی که در زمینه بهبود عملکرد در دهه آخر قرن بیستم به وقوع پیوست، موضوع شناخته شدن و قابلیت اندازه‌گیری میزان رضایت مشتری به مثابه یکی از عناصر و الزامات اصلی نظام‌های مدیریتی در مؤسسات و بنگاه‌های اقتصادی کسب و کار بود. پیاده‌سازی و اجرای سیستم‌های اندازه‌گیری رضایت مشتری، یکی از مهمترین شاخص‌های بهبود عملکرد در سازمان‌های امروزی به شمار می‌رود. امروزه و در میان انبوه رسانه‌های اجتماعی، محتویات تولید شده توسط کاربران و اطلاعات جغرافیایی ارائه شده توسط شهروندان نوع جدیدی از جمع‌آوری داده‌ها را ایجاد کرده است. برای مثال، محتویات تولید شده توسط کاربران روش‌های جدیدی را برای درک بهتر رفتار سفر در مقیاس‌های مختلف فضایی و زمانی ارائه می‌دهد در عین حال، رضایت مشتری به عنوان یکی از شناخته شده‌ترین شاخص‌های ارزیابی مشتری است که توسط محققان مورد استفاده قرار می‌گیرد و شرکت‌های مختلف جهت ارتقاء سطح کیفی خدمات خود از آن استفاده می‌کنند. همچنین با توجه به مطالعات صورت گرفته در حوزه عملکرد محصول جدید محتویات تولید شده توسط کاربران، اطلاعات مهمی در مورد برند، محصولات و عملکرد سهام شرکت‌ها در دسترس قرار می‌دهد. طی دهه‌های گذشته، استفاده از مدل‌های فضایی به یک حوزه تحقیقاتی جدید در بازاریابی تبدیل شده است. **بردلو و همکاران** نشان دادند که همبستگی فضایی می‌تواند بر میزان رضایت‌مندی مشتریان از یک خدمت تاثیر زیادی داشته باشد. آنها روشی برای درک فرایندهای تصمیم‌گیری و یافتن همبستگی فضایی در بحث رضایت‌مندی مشتری معرفی نمودند. با استفاده از روش آنها قصد داریم مطالعه‌ای اکتشافی در خصوص اهمیت جغرافیایی میزان رضایت مشتری از صنعت هتل‌داری را انجام دهیم. این مطالعه بر اساس داده‌های به دست آمده از برنامه TripAdvisor.com در شهر سانفرانسیسکو واقع در ایالت کالیفرنیا آمریکا می‌باشد.

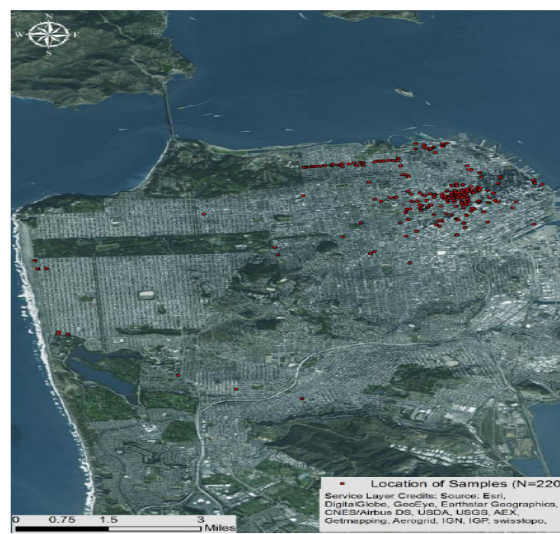
۲ معرفی داده‌ها

در این مطالعه تعداد ۱۷۲۰۵۰ نظر ثبت شده مشتریان در مورد ۲۲۳ هتل در شهر سانفرانسیسکو ایالت کالیفرنیا از تاریخ ۲۲/۰۹/۲۰۰۲ تا ۳۰/۰۳/۲۰۱۵ از طریق TripAdvisor.com جمع‌آوری شده است. این سایت نه تنها یک ابزار پیشگام در تولید محتوا توسط مشتریان است بلکه به عنوان یکی از بزرگترین سایت‌های سیر و سفر در جهان با بیش از ۶۰ میلیون کاربر و بالغ بر ۱۷۰ میلیون بازدیدکننده محسوب می‌شود که فعالیت هتل‌ها، رستوران‌ها، جاذبه‌ها و سایر کسب و کارهای مرتبط در حوزه سرگرمی را در سراسر جهان رصد می‌کند. این سایت پس از بررسی آدرس ایمیل کاربر بر اساس فعالیت‌های او در سایت‌های گردشگری و هتل‌ها اطلاعات آماری و نحوه جستجوی کاربر را جمع‌آوری می‌کند، سپس این اطلاعات در قالب یک گزارش توسط تیمی از متخصصان تحلیل داده ارزیابی می‌شود. نتایج این بررسی‌ها به مالک یا مدیر موسساتی که در فهرست این سایت قرار دارد بصورت پیام ارسال می‌شود. این ویژگی‌ها منجر به افزایش کیفیت محتویات بررسی شده در مقایسه با سایر منابع تولید محتوا و نظر کاربران می‌شود. داده‌های این مطالعه توسط یک تیم بازاریابی با هدف بررسی میزان رضایت مشتریان از خدمات هتل‌ها جمع‌آوری شده است. سپس بر اساس نظر **ژانگ و همکاران** هتل‌هایی که کاربری برای آنها نظری ثبت نکرده و یا نظرات کاربران ناقص ثبت شده است از مجموعه داده‌ها حذف شده‌اند. در نهایت، برای کنترل اختلافات در فرهنگ بین‌الملل و زمینه‌های سیاسی، بررسی‌های ارائه شده توسط کاربران بین‌المللی حذف شده و فقط نظرات کاربران بومی آن منطقه در دسته‌بندی‌ها قرار گرفته است. در نهایت مجموعه نهایی داده‌ها شامل ۷۷۰۹۸ نظر ارائه شده توسط کاربران ایالات متحده در مورد ۲۲۰ هتل جمع‌آوری شده است که موقعیت مکانی هتل‌ها در شکل ۱ آورده شده

است.

۱.۲ همبستگی فضایی رضایت‌مندی مشتریان

در اغلب روش‌های آماری فرض بر این است که مشاهدات تحت شرایط یکسان و به صورت مستقل از هم جمع‌آوری شده‌اند. اما اگر مشاهدات مستقل نباشند و وابستگی آنها تابعی از فاصله بین موقعیت‌های مشاهدات باشد، به‌گونه‌ای که مشاهدات نزدیک به هم وابسته‌تر و مشاهدات دورتر از هم وابستگی کمتری داشته باشند، پس مشاهدات مربوطه، داده‌های فضایی نامیده می‌شوند. **محمدزاده** مجموعه روش‌های تحلیل داده‌های فضایی را آمار فضایی نامیده است. برای تحلیل داده‌های فضایی لازم است یک مدل آماری در نظر گرفته شود. در آمار فضایی به‌طور معمول یک میدان تصادفی به‌عنوان مدل آماری داده‌های فضایی در نظر گرفته می‌شود.



شکل ۱: نظرات کاربران در خصوص کیفیت خدمات هتل‌ها بر اساس توزیع جغرافیایی ۲۲۰ هتل در سانفرانسیسکو

میدان تصادفی مجموعه‌ای از متغیرهای تصادفی مانند $\{Z(s); s \in D\}$ است، که در آن مجموعه اندیس‌گذار D زیر مجموعه‌ای از فضای اقلیدسی d بعدی، $d \geq 1$ ، از R^d است.

برای میدان تصادفی $Z(\cdot)$ ، میانگین در موقعیت s و کوواریانس در موقعیت‌های s_1 و s_2 به‌ترتیب به صورت

$$\begin{aligned}\mu(s) &= E[Z(s)], \quad s \in D \\ C(s_1, s_2) &= Cov[Z(s_1), Z(s_2)], \quad s_1, s_2 \in D\end{aligned}$$

تعریف می‌شوند. برای $s_1 = s_2 = s$ ، واریانس میدان تصادفی $Z(\cdot)$ در مکان s به صورت

$$Var[Z(s)] = E[Z(s) - \mu(s)]^2$$

حاصل می‌شود.

واریانس و کوواریانس در آمار کلاسیک به ترتیب میزان پراکندگی و شباهت یا همبستگی داده‌ها را نشان می‌دهند. در آمار فضایی نیز مفاهیمی تعریف شده‌اند که مشابه واریانس و کوواریانس بیانگر ساختار همبستگی فضایی داده‌ها باشند.

^۱Variogram

جدول ۱: آمار توصیفی نظرات مشتریان

تعداد نمونه	میانگین	انحراف معیار	مینیمم	ماکسیمم
۲۲۰	۲۲/۰	۰۹/۰	-۳۷/۰	۵/۰
۲۲۰	۶۱/۳	۷۵/۰	۱	۵

واریانس تفاضل میدان تصادفی در دو موقعیت s و $s + h$ را تغییرنگار^۱ می‌نامند که به صورت

$$\gamma(h) = \text{Var}(Z(s+h) - Z(s))$$

تعریف می‌شود. تابع $\gamma(h)$ نیز نیم‌تغییرنگار^۲ نامیده می‌شود. کوواریانس $Z(s)$ و $Z(s+h)$ ، یعنی

$$C(h) = \text{Cov}(Z(s), Z(s+h))$$

هم‌تغییرنگار^۳ است و با شرط $C(0) > 0$ عبارت

$$\rho(s, s+h) = \rho(h) = \frac{C(h)}{C(0)}$$

همبستگی‌نگار^۴ نامیده می‌شود. برای تعیین ساختار همبستگی فضایی داده‌ها می‌توان از تغییرنگار، هم‌تغییرنگار یا همبستگی‌نگار استفاده کرد. برآورد تجربی یا کلاسیک هم‌تغییرنگار به صورت

$$\hat{C}(h) = \frac{1}{N_h} \sum_{N(h)} (Z(s_i) - \bar{Z})(Z(s_j) - \bar{Z})$$

تعریف می‌شود، که در آن N_h تعداد زوج نقاطی است که فاصله بین آنها برابر h است.

در داده‌های مربوط به میزان رضایت‌مندی مشتریان نیز در برخی مواقع همبستگی فضایی وجود دارد. در مطالعه حاضر از تجزیه و تحلیل احساسات مشتری (NLP) که روشی بر اساس یادگیری ماشین است استفاده می‌شود به این صورت که بر اساس نظر مشتریان و همچنین جهت‌گیری مشتری که به عنوان یک متغیر کمی تعریف می‌شود عددی پیوسته در بازه -1 تا 1 نسبت داده می‌شود که در آن نقطه صفر وضعیت خنثی را نشان می‌دهد. در عین حال، جهت‌گیری مثبت نشان‌دهنده احساس مثبت و جهت‌گیری منفی احساس منفی مشتری را بیان می‌کند. علاوه بر این، میزان رضایت‌مندی به صورت یک مقیاس لیکرت پنج نقطه‌ای نیز اندازه‌گیری می‌شود، که در آن عدد یک نشان‌دهنده کمترین میزان رضایت و عدد پنج به معنای بیشترین میزان رضایت است. آمارهای توصیفی مربوط به این دو متغیر اندازه‌گیری شده در جدول ۱ نشان داده شده است. گستره متغیر جهت‌گیری این مطالعه از -0.37 تا 0.5 با مقدار میانگین 0.22 است؛ در حالی که نمره عددی رضایت در مقیاس لیکرت از 1 تا 5 تغییر می‌کند که مقدار میانگین 3.61 است. اطلاعات بدست آمده نشان‌دهنده ارتباط بسیار بالایی بین جهت‌گیری ضمنی و رتبه عددی میزان رضایت اندازه‌گیری شده است که توسط ضریب همبستگی گشتاوری پیرسون محاسبه شده و مقدار آن $R = 0.86$ است.

در حالت استقلال مشاهدات استفاده از روش حداقل مربعات به عنوان یک روش برای برآورد پارمترهای رگرسیونی شناخته می‌شود. در این حالت به منظور حفظ ویژگی و اعتبار نتایج آماری مربوط به جمعیت تشکیل‌دهنده، لازم است فرضیات زیر در نظر گرفته شود.

² Semi variogram

³ Covariogram

⁴ Correlogram

۱ - میانگین خطا صفر است.

۲ - خطاها غیرهمبسته با واریانس ثابت هستند.

۳ - خطا دارای توزیع نرمال است.

با این حال، در بسیاری از مواقع در مطالعات بازاریابی، مقدار مشاهده شده در یک مکان اغلب به مقادیر در مکان‌های همسایه وابسته است. یعنی میزان رضایت‌مندی مشتریان از خدمات هتل‌ها تا حد زیادی وابسته به محلی است که هتل در آن قرار دارد و عموماً به خاطر مسایل اقتصادی و اجتماعی هتل‌هایی که نزدیک به هم در یک محل خاص قرار می‌گیرند کیفیت خدماتشان بسیار به هم شباهت دارد. لذا در این گونه بررسی‌های رضایت‌سنجی، داده‌ها مستقل نبوده و دارای همبستگی فضایی خواهند بود در ادامه با معرفی ماتریس‌های مختلف نحوه در نظر گرفتن همسایگی ناحیه‌ها برای لحاظ کردن همبستگی فضایی مورد بررسی قرار می‌گیرد.

نحوه ورود همبستگی فضایی در هنگام تحلیل داده‌های ناحیه‌ای از طریق ماتریس همسایگی با مجاورت^۵ در مدل منظور می‌شود. اما برای وارد کردن اثر فضایی در مدل، ماتریس‌های همسایگی مستقیماً استفاده نمی‌شوند بلکه ابتدا ماتریس وزن فضایی براساس ماتریس همسایگی ساخته شده، سپس از ماتریس وزن برای ورود اثرات فضایی به مدل استفاده می‌شود. در ادامه انواع ماتریس‌های وزن فضایی معرفی خواهد شد. به‌طور کلی ماتریس وزن براساس ماتریس همسایگی ساخته می‌شود که میزان نزدیکی یا به عبارت دیگر تأثیرگذاری دو موقعیت یا ناحیه بر یکدیگر را نشان می‌دهد. با استاندارد کردن سطری ماتریس همسایگی، ماتریس وزن حاصل می‌شود به‌طوری که جمع هر سطر ماتریس وزن برابر یک است. در ادامه نمونه‌هایی از ماتریس همسایگی معرفی می‌شوند.

ماتریس وزن مرز مشترک: برای ساخت این ماتریس ابتدا مؤلفه‌های قطر اصلی ماتریس برابر صفر در نظر گرفته می‌شوند. سپس برای نواحی‌ای که مرز مشترک دارند عدد یک و نواحی‌ای که مرز مشترک ندارند عدد صفر منظور می‌شود. آنگاه تمامی مؤلفه‌های هر سطر بر مجموع مؤلفه‌های همان سطر تقسیم می‌شوند تا ماتریس وزن حاصل شود. به عبارت دیگر عضو z_{ij} ماتریس همسایگی به‌صورت

$$C_{ij} = \begin{cases} 1, & j \in N_1(i) \\ 0, & j \notin N_1(i) \end{cases} \quad (1.2)$$

تعریف می‌شود، که در آن $N_1(i)$ مجموعه نواحی مجاور ناحیه i ام است و $j \in N_1(i)$ به معنی مجاور بودن ناحیه j با ناحیه i است. به‌طور مثال برای یک شبکه ۳ در ۳ به‌صورت ماتریس‌های همسایگی C و وزن W به ترتیب عبارتند

۳	۲	۱
۶	۵	۴
۹	۸	۷

^۵Adjacency matrix

از

$$C = \begin{pmatrix} 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 0 \\ 0 & 1 & 0 & 1 & 0 & 1 & 0 & 1 & 0 \\ 0 & 0 & 1 & 0 & 1 & 0 & 0 & 0 & 1 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 1 & 0 \end{pmatrix} \quad W = \begin{pmatrix} 0 & \frac{1}{2} & 0 & \frac{1}{2} & 0 & 0 & 0 & 0 & 0 \\ \frac{1}{3} & 0 & \frac{1}{3} & 0 & \frac{1}{3} & 0 & 0 & 0 & 0 \\ 0 & \frac{1}{2} & 0 & 0 & 0 & \frac{1}{2} & 0 & 0 & 0 \\ \frac{1}{3} & 0 & 0 & 0 & \frac{1}{3} & 0 & \frac{1}{3} & 0 & 0 \\ 0 & \frac{1}{4} & 0 & \frac{1}{4} & 0 & \frac{1}{4} & 0 & \frac{1}{4} & 0 \\ 0 & 0 & \frac{1}{3} & 0 & \frac{1}{3} & 0 & 0 & 0 & \frac{1}{3} \\ 0 & 0 & 0 & \frac{1}{2} & 0 & 0 & 0 & \frac{1}{2} & 0 \\ 0 & 0 & 0 & 0 & \frac{1}{3} & 0 & \frac{1}{3} & 0 & \frac{1}{3} \\ 0 & 0 & 0 & 0 & 0 & \frac{1}{2} & 0 & \frac{1}{2} & 0 \end{pmatrix}$$

اگر علاوه بر همسایگی‌های با مرز مشترک، همسایگی‌های با k گام فاصله از هر ناحیه نیز به‌عنوان همسایه ناحیه i در نظر گرفته شوند، عناصر ماتریس همسایگی k ناحیه نزدیک به صورت

$$C_{ij} = \begin{cases} 1, & j \in N_k(i) \\ 0, & j \notin N_k(i) \end{cases}$$

تعریف می‌شوند، که در آن $N_k(i)$ نشان‌دهنده k ناحیه مجاور ناحیه i است. همچنین می‌توان ماتریس همسایگی را براساس همبستگی میان نقاط به صورت

$$C_{ij} = r_{ij}^{-b}$$

بیان کرد، که در آن b ضریب اصطکاک فاصله^۶ و r_{ij} همبستگی میان دو ناحیه i و j است. ضریب اصطکاک b معمولاً برابر ۱ در نظر گرفته می‌شود. **کلیف و ارد** از این ساختار برای ساخت ماتریس وزن فضایی استفاده کردند، به همین علت ماتریس وزن حاصل از این ماتریس همسایگی، **کلیف و ارد**^۷ نامیده می‌شود. یکی دیگر از ماتریس‌های همسایگی برحسب همبستگی میان نواحی به صورت

$$C_{ij} = \begin{cases} 1, & r_{ij} \leq r_0 \\ 0, & r_{ij} > r_0 \end{cases}$$

است، که در آن r_0 یک آستانه برای میزان همبستگی است. ماتریس وزن براساس همبستگی شبه عمومی^۸ نیز به صورت

$$C_{ij} = \exp\left(-\frac{r_{ij}}{\bar{r}}\right)$$

تعریف می‌شود. یکی دیگر از راه‌های شناسایی میزان همبستگی فضایی موقعیت‌ها، فاصله میان موقعیت‌ها است. ماتریس همسایگی بر اساس حد فاصل میان نواحی به صورت

$$C_{ij} = \begin{cases} 1, & 0 \leq d_{ij} \leq d \\ 0, & i = j \text{ یا } d_{ij} > d \end{cases}$$

^۶Distance friction coefficient

^۷Cliff-Ord

^۸Quasi-global correlation

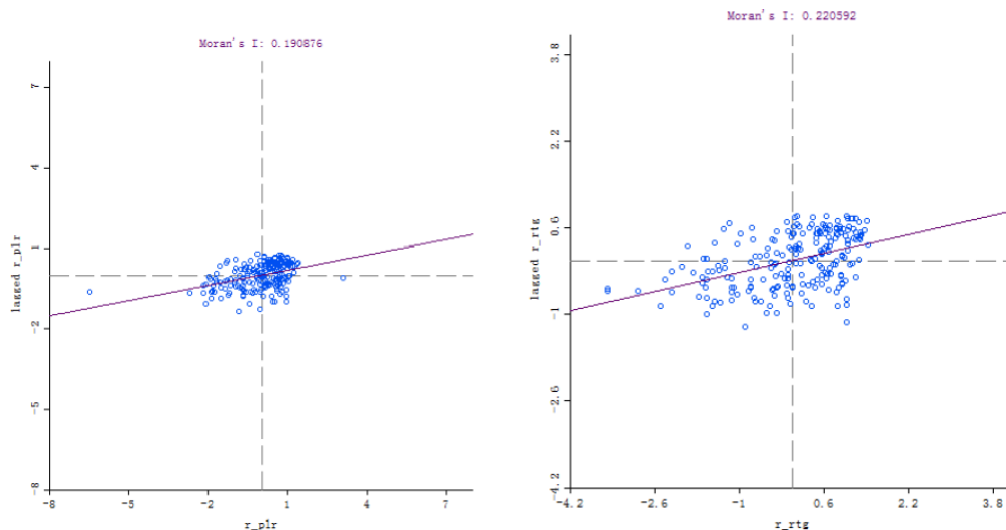
تعریف می‌شود، که در آن d_{ij} فاصله بین دو ناحیه i و j است و d آستانه است. چند نمونه از ماتریس‌های وزن که براساس فاصله بین نواحی حاصل می‌شوند به صورت

$$C_{ij} = \begin{cases} 1, & i = j \\ d_{ij}^{-\alpha}, & i \neq j \end{cases} \quad (۲.۲)$$

$$C_{ij} = \begin{cases} 1, & i = j \\ \exp(-\beta d_{ij}^{\alpha}), & i \neq j \end{cases} \quad (۳.۲)$$

$$C_{ij} = \begin{cases} [1 - (\frac{d_{ij}}{d})^k]^k, & 0 \leq d_{ij} \leq d \\ 1, & i = j \text{ یا } d_{ij} > d \end{cases} \quad (۴.۲)$$

هستند، که در آن‌ها α ، β و k اعدادی مثبت هستند. برای ماتریس همسایگی (۲.۲) مقدار α معمولاً برابر ۱ یا ۲ در نظر گرفته می‌شود.



شکل ۲: نتایج آزمون موران

پیش از تصمیم‌گیری در مورد وابستگی فضایی، لازم است وجود وابستگی فضایی را تایید کنیم، با توجه به آنکه مسئله مهم برای تحلیل رضایت مشتری تشخیص وجود خود همبستگی فضایی^۹ است. برای این منظور معمولاً از آزمون موران^{۱۰} استفاده می‌شود. فرض صفر این آزمون عدم وجود خود همبستگی فضایی در داده‌ها است. آماره این آزمون براساس وزن‌های بین نواحی به صورت

$$I = \frac{n \sum_{i=1}^n \sum_{j=1}^n w_{ij} (t_i - \bar{t})(t_j - \bar{t})}{(\sum_{i=1}^n \sum_{j=1}^n w_{ij})(\sum_{i=1}^n (t_i - \bar{t})^2)}$$

است، که در آن n تعداد داده‌ها، w_{ij} وزن میان دو ناحیه i و j مقدار ویژگی مورد بررسی در ناحیه i است. وقتی

^۹Spatial autocorrelation

^{۱۰}Moran's I test

$|I^*| > z_{1-\alpha/2}$ ، فرض وجود خود همبستگی فضایی داده‌ها پذیرفته می‌شود، که در آن

$$I^* = \frac{I - E(I)}{\sqrt{Var(I)}}, \quad P(Z \leq 1 - \alpha/2) = z_{1-\alpha/2}$$

که در آن Z متغیر تصادفی با توزیع نرمال استاندارد است. با توجه به فرض تأثیرگذاری موقعیت قرارگیری هتل‌ها بر رضایت مشتریان از آزمون موران انجام شد. با در نظر گرفتن ماتریس مجاورت (۱.۲) فرض صفر آزمون موران در سطح معنی‌داری ۰/۰۵ با p -مقدار کمتر از ۰/۰۰۱ رد می‌شود و فرضیه وجود خود همبستگی فضایی در داده‌ها تایید می‌گردد.

بحث و نتیجه‌گیری

نتایج موران برای هر دو متغیر مورد بررسی یعنی جهت‌گیری مشتری و رضایت‌مندی مشتری در شکل ۲ رسم شده است. یک همبستگی فضایی مثبت برای رضایت مشتریان هتل‌ها در منطقه مورد مطالعه وجود دارد. بر اساس آماره موران، هتل‌هایی که جهت‌گیری مشتری بالایی دارند، از نظر فضایی در گروه هتل‌های با رضایت بالای مشتری قرار می‌گیرند. بنابراین، رویکرد مکان‌یابی هتل بسیار مهم است. مانند هر شرکت دیگری هتل‌ها نیز برای به حداکثر رساندن میزان سود خالص خود در رقابت با سایرین با استفاده از تحلیل‌های مکان‌یابی و اقدامات پیش‌بینی شده رقبا مکان خود را تعیین می‌کنند. نکته مهم، چانگ (۲۰۰۷) از این یافته‌ها هنگام استفاده از رضایت مشتری به عنوان متغیر وابسته برای تجزیه و تحلیل تجربی پیش‌آمدها و پیامدها، نشان می‌دهد که وجود وابستگی فضایی به طور بالقوه فرض نااریب بودن برآورد پارامترهای مدل رگرسیونی را نقض می‌کند. لذا برای دستیابی به نتایج مطلوب‌تر لازم است با وارد کردن همبستگی فضایی و استفاده از رویکرد فضایی تحلیل بهتری از میزان رضایت‌مندی مشتریان از خدمات هتل‌ها صورت گیرد. رضایت مشتری توسط محققان به طور گسترده مطالعه شده و مورد استفاده شرکت‌ها قرار گرفته است. گسترش محتویات تولید شده توسط کاربران و اطلاعات موقعیت جغرافیایی به مشتریان و شرکت‌ها در راستای رسیدن به کیفیت مطلوب‌تر کمک می‌کند. در این مقاله، با توجه به اندازه‌گیری آماره موران، نشان دادیم که هتل‌هایی با رتبه بالا در یک دسته قرار می‌گیرند و مکان‌یابی برای آنها اهمیت دارد. در مطالعات آتی قصد داریم با تعیین یک ماتریس همسایگی مناسب و استفاده از مدل‌های اتورگرسیون شرطی به پیش‌گویی میزان رضایت‌مندی مشتریان در یک مکان خاص که قرار است هتلی در آنجا احداث شود بپردازیم. و در نهایت با استفاده از این روش بهترین نقطه مکانی برای احداث هتل به سرمایه‌گذاران این صنعت معرفی شود.

تقدیر و تشکر

از استادان ارجمندم جناب آقای دکتر محسن محمدزاده، آقای دکتر موسی گلعلی‌زاده، آقای دکتر کیومرث مترجم بسیار سپاسگزارم چرا که بدون راهنمایی‌ها و آموزش‌های ایشان تهیه این مقاله بسیار مشکل می‌نمود.

مراجع

محمدزاده، م.، (۱۳۹۸)، *آمار فضایی و کاربردهای آن*، چاپ سوم، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران،

Alcácer, J., Chung, W. (2007), *Location Strategies and Knowledge Spillovers*, Management Science 53, 760-776.

Bradlow(2005), *Spatial models in marketing*, Marketing Letters 16, 267-278.

Cliff, A. D. and Ord, J. K. (1973), *Spatial Autocorrelation*, Pion Limited, London.

Moran, P. A. P. (1950), *Notes on Continuous Stochastic Phenomena*, **37**, 17-23.

Wang(2015), *Crowdsourcing the Landscape of Cannabis (Marijuana) of the Contiguous United States*, Environment and Planning A 10.1177/0308518x15598541.

Zhang(2015), *Social Learning in Networks of Friends versus Strangers*, Marketing Science, **34**, 573-589.



تحلیل کواریانس گاوسی وارون برای داده‌های فضایی

زهرا زارع‌پور یکنمی، افشین فلاح^۱
گروه آمار، دانشگاه بین‌المللی امام خمینی

چکیده: نرمال بودن توزیع متغیر پاسخ و مستقل بودن مشاهدات در هر تیمار، از جمله فرض‌های پایه‌ای تحلیل کواریانس به‌شیوه‌ی سنتی هستند، به‌گونه‌ای که برقرار نبودن این فرض‌ها اعتبار نتایج حاصل از این روش را مخدوش می‌کند. در این مقاله برای تحلیل کواریانس در جوامعی که مشاهدات مثبت و دارای وابستگی فضایی هستند، مدلی پیشنهاد شده است. چگونگی برآزش مدل پیشنهادی از دیدگاه بیزی مورد بحث قرار گرفته و از یک مطالعه‌ی شبیه‌سازی برای ارزیابی آن استفاده شده است.

واژه‌های کلیدی: تحلیل کواریانس، گاوسی وارون، داده‌های فضایی، رهیافت بیزی.
کد موضوع‌بندی ریاضی (۲۰۱۰): 62H11, 62J10.

۱ مقدمه

رگرسیون و تحلیل واریانس از جمله مهم‌ترین روش‌های مرسوم برای بیان روابط بین متغیرها هستند. اما در برخی کاربردها هیچ یک از این روش‌های مدل‌سازی به‌تنهایی نمی‌تواند پاسخ مناسب به پرسش‌های پژوهش‌گر را فراهم آورند. تحلیل کواریانس ابزار نیرومندی است که از ترکیب دو روش رگرسیون و تحلیل واریانس ساخته می‌شود و از مزایای هر دو روش استفاده می‌کند. تحلیل کواریانس نخستین بار توسط **فیشر (۱۹۳۲)** مطرح شده است. **کاکس و مک کالک (۱۹۸۲)** برخی از جنبه‌های تحلیل کواریانس در زیست‌سنجی را مورد بررسی قرار دادند. روش‌های بیزی و غیربیزی در رابطه با مدل‌های تحلیل کواریانس تحت شرایط ناهم‌واریانسی، توسط **آناندا (۱۹۹۸)** مورد مطالعه قرار گرفته است. **میلر و چپمن (۲۰۰۱)** پذیره‌های زیربنایی تحلیل کواریانس را مورد بحث قرار داده‌اند. از جمله مهم‌ترین پذیره‌های زیربنایی تحلیل کواریانس مرسوم نرمال بودن توزیع متغیر پاسخ و مستقل بودن مشاهدات مربوط به تیمارهای گوناگون است. این در حالی است که در بسیاری از کاربردها، این پذیره‌ها برقرار نیستند. در این راستا، **مشکانی و همکاران (۲۰۱۴)** تحلیل کواریانس گاوسی وارون را برای استفاده در شرایطی که مشاهدات پاسخ نامنفی و چوله به راست هستند، توسعه داده‌اند. **فلاح (۲۰۱۶)**، تحلیل

^۱ نام و ایمیل ارائه دهنده مقاله: افشین فلاح، afshin.ikiu@gmail.com

کواریانس گاوسی آمیخته‌ی متناهی را برای شرایطی که جامعه تحت مطالعه ناهمگن بوده و از زیرجامعه‌هایی با ویژگی‌های متفاوت تشکیل شده، مورد بررسی قرار داده است. از طرفی در برخی از کاربردها مشاهدات مستقل نیستند و به دلیل قرار گرفتن در موقعیت‌های فضایی خاص خود دارای وابستگی فضایی هستند. در این شرایط استفاده از نظریه‌ی مرسوم تحلیل کواریانس به دلیل لحاظ نکردن وابستگی‌های فضایی بین مشاهدات به نتایج گمراه کننده می‌انجامد.

در این مقاله، تحلیل کواریانس گاوسی وارون برای داده‌های فضایی از دیدگاه بیزی مورد بررسی قرار گرفته و مدلی برای آن پیشنهاد شده است. در بخش ۲، چارچوب نظری مدل پیشنهادی شرح داده شده است. در بخش ۳، نحوه برازش مدل از طریق رهیافت بیزی مورد بحث قرار گرفته است. در بخش ۴، برای ارزیابی مدل پیشنهادی از یک مطالعه شبیه‌سازی استفاده شده است.

۲ مدل پیشنهادی

به پیروی از [مشکانی و همکاران \(۲۰۱۴\)](#)، مدل تحلیل کواریانس را می‌توان به صورت

$$\mu^* = (\mu_{11}^{-1}, \dots, \mu_{tr}^{-1}) = X\beta + Z\gamma = W\theta,$$

در نظر گرفت، که در آن μ بردار میانگین مشاهدات پاسخ، X ماتریس طرح، Z ماتریس متغیرهای کمکی، β بردار ضرایب تیماری، γ بردار ضرایب رگرسیونی، $W = [X|Z]$ ماتریس مشاهدات و $\theta = (\beta', \gamma')'$ بردار ضرایب مدل را نشان می‌دهند. فرض کنید متغیرهای پاسخ در مدل تحلیل کواریانس دارای توزیع گاوسی وارون به صورت

$$y_{ij}|X, Z \sim f_{IG}(y_{ij}|\mu_{ij}, \lambda), \quad i = 1, \dots, t; j = 1, \dots, r,$$

هستند، که در آن

$$f_{IG}(y_{ij} | \mu_{ij}, \lambda) = \sqrt{\frac{\lambda}{2\pi y_{ij}^3}} \exp \left\{ -\frac{\lambda(y_{ij} - \mu_{ij})^2}{2y_{ij}\mu_{ij}^2} \right\}, \quad y_{ij}, \mu_{ij}, \lambda > 0,$$

تابع چگالی توزیع گاوسی وارون، μ_{ij} میانگین متغیر پاسخ برای مشاهده مربوط به تکرار j -ام از تیمار i -ام و λ پارامتر مقیاس را نشان می‌دهند. در مدل پیشنهادی اثر وابستگی‌های فضایی موجود بین مشاهدات پاسخ در هر تیمار، به وسیله‌ی بردار اثرهای تصادفی فضایی $\phi = (\phi_{11}, \dots, \phi_{tr})'$ در تابع پیوند مدل، لحاظ می‌شوند. فرض کنید بردار مشاهدات پاسخ در موقعیت‌های مکانی $(s_{11}, \dots, s_{tr})$ به صورت $\{y_{ij} = y(s_{ij}); D \subset \mathbb{R}^d, d \geq 2\}$ تعریف می‌شوند. در این صورت، تابع پیوند مدل تحلیل کواریانس گاوسی وارون فضایی به صورت

$$\mu_{ij}^{-1} = w_{ij}\theta + \phi_{ij} \quad (1.2)$$

خواهد بود، که در آن $w_{ij} = (x_i, z_{ij})$ سطر ij -ام ماتریس W و ϕ_{ij} اثر تصادفی فضایی متناظر با مشاهده مربوط به تکرار j -ام از تیمار i -ام است. اثر تصادفی فضایی ϕ_{ij} را می‌توان به صورت شرطی بر اثرهای تصادفی فضایی مشاهدات همسایه در تیمار i -ام به صورت

$$\phi_{ij}|\phi_{i(-j)}, \sigma^2 \sim N \left(\xi \sum_{l \neq j} \frac{a_{ijl}}{a_{ij+}} \phi_{il}, \frac{\sigma^2}{a_{ij+}} \right), \quad (2.2)$$

در نظر گرفت، که در آن $\phi_{i(-j)} = \{\phi_{il}; l \neq j\}$ ، واریانس مشترک اثرهای تصادفی فضایی، a_{ijl} درایه (j, l) از ماتریس همسایگی است که به صورت

$$a_{ijl} = \begin{cases} 1, & \text{اگر مشاهده‌ی } i \text{ در تکرارهای } j \text{ و } l \text{ همسایه باشند.} \\ 0, & j = l \end{cases} \quad (3.2)$$

تعریف می‌شود، $a_{ijl} = \sum_{l \neq j} a_{ijl}$ و ξ پارامتر هموارسازی فضایی است. پارامتر هموارسازی در بازه‌ی $\xi \in \left(\frac{1}{\lambda_{(1)}}, \frac{1}{\lambda_{(n)}}\right)$ تغییر می‌کند که در آن $\lambda_{(1)} \geq \dots \geq \lambda_{(n)}$ مقادیر ویژه ماتریس $D_M^{-\frac{1}{\xi}} M D_M^{-\frac{1}{\xi}}$ هستند، $M = (m_{ijl}) = \left(\frac{a_{ijl}}{a_{ij+}}\right)$ ماتریس وزن و D_M ماتریس قطری با درایه‌های قطری $\{m_{11+}, \dots, m_{tr+}\}$ است (کاوسی و همکاران، ۲۰۰۹). با توجه به اینکه همبستگی‌های فضایی تنها بین مشاهدات متناظر با تکرارهای گوناگون در هر تیمار وجود دارند، ماتریس همسایگی در مدل پیشنهادی، ماتریسی بلوکی و به صورت

$$A = \begin{bmatrix} A_1 & \dots & 0 \\ \vdots & \ddots & \vdots \\ 0 & \dots & A_t \end{bmatrix}, \quad (4.2)$$

در نظر گرفته می‌شود، که در آن ماتریس A_i مشاهدات همسایه تیمار i -ام، به‌ازای $i = 1, \dots, t$ ، را نشان می‌دهند.

۳ برآورد پارامترهای مدل

توابع درستنمایی و تابع لگ درستنمایی متناظر با مدل تحلیل کواریانس پیشنهادی را می‌توان به ترتیب به صورت

$$L(\Psi|y, W) = \prod_{i=1}^t \prod_{j=1}^r (f_{IG}(y_{ij}|\mu_{ij}, \lambda)),$$

$$\ell(\Psi|y, W) = \sum_{i=1}^t \sum_{j=1}^r (f_{IG}(y_{ij}|\mu_{ij}, \lambda)),$$

نوشت. این توابع را می‌توان با استفاده از نمادهای برداری و ماتریسی، به ترتیب به صورت

$$\ell(\Psi|y, W) = \sum_{i=1}^t \sum_{j=1}^r \log(\lambda) - \sum_{i=1}^t \sum_{j=1}^r \log(\pi y_{ij}^{\frac{1}{\lambda}}) - \frac{\lambda}{\gamma} (Q_1(\theta_1) + Q_5(\theta, \phi)),$$

$$L(\Psi|y, W) = \prod_{i=1}^t \prod_{j=1}^r \left(\frac{\lambda}{\pi y_{ij}^{\frac{1}{\lambda}}} \right) \exp \left\{ \frac{\lambda}{\gamma} (Q_1(\theta) + Q_4(\theta, \phi)) \right\}, \quad (1.3)$$

بازنویسی نمود، که در آن

$$Q_1(\theta) = \theta' \left(\sum_{i=1}^t \sum_{j=1}^r y_{ij} w'_{ij} w_{ij} \right) \theta_k - \gamma \left(\sum_{i=1}^t \sum_{j=1}^r w_{ij} \right) \theta + \sum_{i=1}^t \sum_{j=1}^r y_{ij}^{-1}$$

$$= (YW\theta - 1_n)' Y^{-1} (YW\theta - 1_n),$$

$$Q_4(\theta, \phi) = \sum_{i=1}^t \sum_{j=1}^r y_{ij} \phi_{ij}^{\frac{1}{\lambda}} + \gamma \sum_{i=1}^t \sum_{j=1}^r y_{ij} \phi_{ij} w_{ij} \theta - \gamma \sum_{i=1}^t \sum_{j=1}^r \phi_{ij}$$

$$= \phi' Y \phi + \gamma \phi' Y W \theta - \gamma 1_n' \phi.$$

در رویکرد بیزی لازم است ابتدا توزیع‌های پیشین مناسبی برای پارامترهای مدل در نظر گرفته شود. برای این منظور، با فرض استقلال پیشینی پارامترهای (θ, λ) و (ϕ, σ^2) ، توزیع پیشین توأم پارامترها را می‌توان به صورت

$$\begin{aligned}\pi(\theta, \phi, \sigma^2, \lambda) &= \pi(\theta, \lambda)\pi(\phi, \sigma^2) \\ &= \pi(\theta|\lambda)\pi(\lambda)\pi(\phi|\sigma^2)\pi(\sigma^2),\end{aligned}\quad (2.3)$$

نوشت. برای پارامترهای θ و λ به پیروی از **مشکانی و همکاران (۲۰۱۶)** توزیع‌های پیشین مزدوج به شکل

$$\begin{aligned}\theta | \lambda &\sim N_p(\eta_*, \lambda^{-1} \Delta_*), \quad \theta, \eta_* \in \mathbb{R}^p, \\ \lambda &\sim G\left(\frac{a_*}{\nu}, \frac{b_*}{\nu}\right) \quad \lambda, a_*, b_* > 0,\end{aligned}\quad (3.3)$$

در نظر گرفته شده‌اند، که در آن‌ها Δ_* یک ماتریس قطری با درایه‌های قطری مثبت $\delta_* = (\delta_{*1}, \dots, \delta_{*p})$ است، $\eta_* = (\eta_{*1}, \dots, \eta_{*p})'$ و $p = t + q$. از این رو، توزیع پیشین توأم (θ, λ) به صورت

$$\begin{aligned}\pi(\theta, \lambda) &\propto \pi(\theta | \lambda)\pi(\lambda) \\ &\propto \lambda^{\frac{(a_*+p)}{\nu}-1} \exp\left\{-\frac{\lambda}{\nu} \left(b_* + (\theta - \eta_*)' \Delta_*^{-1} (\theta - \eta_*)\right)\right\},\end{aligned}\quad (4.3)$$

محاسبه می‌شود و یک توزیع گاما-نرمال است. برای بردار اثرهای تصادفی $\phi = (\phi_{11}, \dots, \phi_{tr})'$ و پارامتر واریانس σ^2 ، به ترتیب توزیع‌های پیشین نرمال چندمتغیره و گامای وارون به صورت

$$\begin{aligned}\phi | \sigma^2 &\sim N_n\left(0, (D_M - \xi M)^{-1} \sigma^2\right), \\ \frac{1}{\sigma^2} &\sim G(f_*, g_*),\end{aligned}\quad (5.3)$$

لحاظ می‌شوند. چگالی‌های متناظر با این توزیع‌های پیشین را می‌توان به شکل متناسب

$$\begin{aligned}\pi(\phi | \sigma^2) &\propto \exp\left\{-\frac{1}{\nu \sigma^2} \phi' (D_M - \xi M)^{-1} \phi\right\}, \\ \pi(\sigma^2) &\propto \sigma^{2-(f_*+1)} \exp\left\{-\frac{g_*}{\sigma^2}\right\},\end{aligned}\quad (6.3)$$

نوشت. در رابطه‌های (۳.۳) و (۵.۳) از اندیس صفر برای نشان دادن ابرپارامترهای مدل بیزی که می‌بایست توسط تحلیل‌گر تعیین شوند، استفاده شده است. با توجه به توزیع‌های پیشین (۲.۳)، (۴.۳) و (۵.۳)، تابع درستنمایی (۱.۳)، توزیع پسین توأم پارامترها به صورت

$$\begin{aligned}\pi(\Psi | y, W) &= L(\Psi | y, W) \times \pi(\theta, \phi, \sigma^2, \lambda) \\ &\propto \prod_{i=1}^t \prod_{j=1}^r \lambda^{\frac{(a_*+p+n)}{\nu}-1} \exp\left\{-\frac{\lambda}{\nu} (Q_\nu(\theta) + Q_\nu(b, \eta_*, \Delta_*) + Q_\nu(\theta, \phi))\right\} \\ &\times \sigma^{2-(f_*+1)} \exp\left\{-\frac{1}{\sigma^2} Q_\delta(\phi)\right\},\end{aligned}\quad (7.3)$$

به دست می‌آید، که در آن

$$\begin{aligned}Q_\nu(\theta) &= (\theta - \Sigma^{-1} d)' \Sigma (\theta - \Sigma^{-1} d), \\ Q_\nu(b, \eta_*, \Delta_*) &= b_* + \eta_*' \Delta_*^{-1} \eta_* + \mathbf{1}_n' Y^{-1} \mathbf{1}_n - d' \Sigma^{-1} d,\end{aligned}$$

$$\begin{aligned}\Sigma &= (\Delta_{\cdot}^{-1} + W'YW), \\ d &= \Delta_{\cdot}^{-1}\eta_{\cdot} + W'\mathbf{1}_n, \\ Q_{\Delta}(\phi) &= \frac{\phi'(DM - \xi M)^{-1}\phi}{2} - g_{\cdot}.\end{aligned}$$

با توجه به پیچیدگی توزیع پسین توأم پارامترها در رابطه‌ی (۷.۳) به‌دست آوردن نمایش بسته برای کمیت‌های پسینی شدنی نیست. از این‌رو، در ادامه از روش‌های مونت کارلو زنجیر مارکوفی برای نمونه‌گیری از این توزیع پسین و تقریب کمیت‌های پسینی مورد علاقه استفاده می‌شود.

۱.۳ توزیع‌های پسین شرطی کامل

در این بخش برای نمونه‌گیری از توزیع پسین، یک الگوریتم گیز توسعه داده می‌شود. برای این منظور لازم است توزیع‌های پسین شرطی کامل متناظر پارامترهای مدل محاسبه شوند. برپایه‌ی توزیع پسین توأم پارامترها در رابطه‌ی (۷.۳)، توزیع پسین شرطی کامل پارامتر λ ، σ^2 ، θ و ϕ به‌ترتیب به‌صورت

$$\begin{aligned}\lambda | others &\sim G\left(\frac{a_{\cdot} + p + n}{2}, \frac{Q_2(\theta) + Q_3(b, \eta_{\cdot}, \Delta_{\cdot}) + Q_4(\theta, \phi)}{2}\right) \\ \frac{1}{\sigma^2} | others &\sim G(f_{\cdot}, Q_{\Delta}(\phi)) \\ \pi(\theta | others) &\propto \prod_{i=1}^t \prod_{j=1}^r \exp\left\{-\frac{\lambda}{2}(Q_2(\theta) + Q_4(\theta, \phi))\right\}, \\ \pi(\phi | others) &\propto \prod_{i=1}^t \prod_{j=1}^r \exp\left\{-\frac{\lambda}{2}Q_4(\theta, \phi) - \frac{1}{\sigma^2}Q_{\Delta}(\phi)\right\},\end{aligned}$$

به‌دست می‌آیند. همان‌گونه که ملاحظه می‌شود این توزیع‌ها نمایش بسته و شناخته شده‌ای ندارند. به‌همین دلیل، برای نمونه‌گیری از آنها لازم است از الگوریتم متروپولیس-هاستینگز استفاده شود. از این‌رو، برای نمونه‌گیری از توزیع پسین توأم پارامترها به یک الگوریتم گیز با تکرارهای آشیانه‌ای متروپولیس-هاستینگز نیاز است. مراحل کلی الگوریتم گیز برای نمونه‌گیری از توزیع پسین (۷.۳) به شرح زیر است:

۱. مقادیر اولیه‌ی $\Psi^{(0)} = \{\theta^{(0)}, \phi^{(0)}, \lambda^{(0)}, \sigma^{2(0)}\}$ را برای پارامترهای مدل در نظر بگیرید.

۲. در تکرار $(m+1)$ -ام الگوریتم، برپایه‌ی مقادیر حاصل از تکرار (m) -ام، پارامترهای مدل را به‌صورت زیر به‌روز کنید:

$$\begin{aligned}\lambda^{(m+1)} | others &\sim G\left(\frac{a_{\cdot} + p + n}{2}, \frac{Q_2(\theta^{(m)}) + Q_3(b, \eta_{\cdot}, \Delta_{\cdot}) + Q_4(\theta^{(m)}, \phi^{(m)})}{2}\right), \\ \frac{1}{\sigma^{2(m+1)}} | others &\sim G\left(f_{\cdot}, Q_{\Delta}(\phi^{(m)})\right), \\ \pi(\theta^{(m+1)} | others) &\propto \prod_{i=1}^t \prod_{j=1}^r \exp\left\{-\frac{\lambda^{(m+1)}}{2}(Q_2(\theta^{(m)}) + Q_4(\theta^{(m)}, \phi^{(m)}))\right\}, \\ \pi(\phi_k^{(m+1)} | others) &\propto \prod_{i=1}^t \prod_{j=1}^r \exp\left\{-\frac{\lambda^{(m+1)}}{2}(Q_4(\theta^{(m+1)}, \phi^{(m)})) - \frac{1}{\sigma^{2(m+1)}}Q_{\Delta}(\phi^{(m)})\right\}.\end{aligned}$$

۳. گام ۲ الگوریتم را تا همگرا شدن زنجیر مارکوف به توزیع مانای متناظر خود تکرار کنید.

۴ مطالعه‌ی شبیه‌سازی

در این بخش برای ارزیابی مدل پیشنهادی از یک مطالعه‌ی شبیه‌سازی استفاده شده است. برای این منظور یک مدل تحلیل کواریانس با یک متغیر کمکی و یک عامل دو سطحی، به‌ازای $(\beta_0 = 4, \beta_1 = 1, \gamma_1 = 2)$ و $\theta = 2$ و $\lambda = 2$ در نظر گرفته شده است. متغیرهای پاسخ y_{ij} ، $i = 1, 2$ و $j = 1, \dots, r$ از توزیع (1.2) با تابع پیوند (1.2) شبیه‌سازی شده‌اند. مقادیر ریشه دوم میانگین توان دوم خطای برآوردگرهای بیزی پارامترهای مدل پیشنهادی محاسبه و در جدول ۱ ارائه شده‌اند. تمامی محاسبات ۴۰۰۰ بار تکرار شده است. همان‌گونه که ملاحظه می‌شود مدل پیشنهادی به‌خوبی وابستگی‌های فضایی مشاهدات را نشان می‌دهد.

جدول ۱: مقادیر ریشه دوم میانگین توان دوم خطا برآوردگرهای ضرایب مدل تحلیل کواریانس گاوسی وارون برای داده‌های فضایی و غیرفضایی در رهیافت بیزی.

اندازه‌ی نمونه			مدل تحلیل کواریانس	پارامتر
۸۰	۴۰	۲۰		
۱/۷۳۵۲	۰/۸۶۶۸	۱/۴۱۰۰	مدل فضایی پیشنهادی	β_0
۲/۰۰۵۴	۱/۰۱۵۵	۳/۹۲۹۳	مدل کلاسیک	
۱/۹۵۷۶	۰/۹۳۱۳	۱/۵۳۳۹	مدل فضایی پیشنهادی	β_1
۲/۰۰۰۶	۱/۰۱۲۳	۳/۸۶۷۵	مدل کلاسیک	
۱/۹۸۴۸	۰/۹۳۲۱	۱/۴۹۲۳	مدل فضایی پیشنهادی	γ
۲/۰۰۴۲	۱/۰۱۳۴	۳/۹۸۰۴	مدل کلاسیک	

بحث و نتیجه‌گیری

تحلیل کواریانس یکی از روش‌های بسیار پرکاربرد مدل‌سازی و تلفیقی از تحلیل رگرسیونی و تحلیل واریانس است. اعتبار نتایج حاصل از این روش، به میزان درستی فرضیات زیر بنایی آن مانند مستقل بودن مشاهدات در هر تیمار و نرمال بودن توزیع مشاهدات پاسخ وابسته است. در شرایطی که این فرضیات برقرار نباشند، نتایج حاصل از تحلیل کواریانس می‌تواند به‌صورت قابل توجه‌ای گمراه‌کننده باشد. هنگامی که مشاهدات پاسخ مثبت و دارای همبستگی فضایی هستند، مدل تحلیل کواریانس مبتنی بر توزیع گاوسی وارون برای داده‌های فضایی که در این مقاله پیشنهاد شده است، می‌تواند به‌عنوان جایگزینی انعطاف‌پذیر برای مدل مرسوم تحلیل کواریانس به‌کار گرفته شود.

مراجع

- Ananda, M. A. (1998), Bayesian and Non-Bayesian Solutions to Analysis of Covariance Models Under Heteroscedasticity, *Statistics for Spatial Data*, **86**(1), 177–192.
- Cox, D. R. and McCullagh, P. (1982), Biometrics Invited Paper with Discussion. Some Aspects of Analysis of Covariance, *Biometrics*, **38**, 541–561.

- Fallah, A. (2016), Gaussian Mixture Analysis of Covariance, *Journal of Statistical Computation and Simulation*, **86**(16), 3158–3174.
- Fisher, R. A. (1932), *Statistical Methods for Research Workers*, 4th Edition, Oliver and Boyd, Edinburgh.
- Kavousi, A., Meshkani, M. R. and Mohammadzadeh, M. (2009), Spatial Analysis of Relative Risk of Lip Cancer in Iran: a Bayesian Approach, *Environmetrics*, **20**, 347–359.
- Meshkani, M. R, Fallah, A. and Kavousi, A. (2014), Analysis of Covariance Under Inverse Gaussian Model, *Journal of Applied Statistics*, **41**(6), 1189-1202.
- Meshkani M. R., Fallah A. and Kavousi A. (2016), Bayesian Analysis of Covariance Under Inverse Gaussian Model, *Journal of Applied Statistics*, **43**, 280–298.
- Miller, G. A. and Chapman, J. P. (2001), Misunderstanding Analysis of Covariance, *Journal of the Abnormal Psychol*, **110**, 40–48.



مدل بندی بیزی داده های نرخ فضایی با مدل رگرسیون بتا فضایی

کبری قلی زاده^۱، محسن محمدزاده

گروه آمار، دانشگاه تربیت مدرس

چکیده: مدل های رگرسیونی بتای فضایی ابزاری مناسب برای مدل بندی داده های نرخ وابسته فضایی است. گاهی در تحلیل این مدل ها علاوه بر مدل بندی پارامتر میانگین نیاز است پارامتر دقت نیز مدل بندی شود، اما به دلیل پیچیده شدن این مدل، استفاده از روش های معمول و الگوریتم های نمونه گیری مونت کارلوی زنجیر مارکوفی موجب افزایش زمان محاسبات می شود. در این مقاله از روش تقریب لاپلاس آشیانی جمع بسته برای تحلیل مدل رگرسیون بتا فضایی استفاده می شود. سپس در یک مطالعه شبیه سازی کارایی مدل ارائه شده مورد ارزیابی قرار گرفته و نحوه کاربست آن در یک مثال واقعی نشان داده شده است.

واژه های کلیدی: مدل رگرسیونی بتا فضایی، تقریب لاپلاس آشیانی جمع بسته.
کد موضوع بندی ریاضی (۲۰۱۰): 91B72، 62M30، 62H11.

۱ مقدمه

در مدل بندی متغیرهای نرخ که در بازه $(0, 1)$ توزیع شده اند، ممکن است استفاده از مدل رگرسیونی با فرض نرمال بودن متغیر پاسخ، پیشگویی هایی برای متغیر پاسخ ارائه دهد که خارج از محدوده $(0, 1)$ باشند. از طرفی استفاده از روش تبدیل برای حل این مساله ممکن است باعث غیرقابل تفسیر شدن پارامترهای حاصل بر حسب متغیر پاسخ واقعی باشد. بنابراین **فراری و کریباری (۲۰۰۴)** مدل رگرسیونی بتا را برای مدل بندی متغیرهایی که دارای توزیع بتا هستند، معرفی کردند. فرض کنید متغیر تصادفی Y دارای توزیع بتا است بطوری که $E(Y) = \mu$ و $Var(Y) = \frac{\mu(1-\mu)}{1+\varphi}$. در این صورت، چگالی احتمال توزیع بتا را می توان بر حسب پارامترهای میانگین μ و دقت φ به صورت

$$f(y|\varphi, \mu) = \frac{\Gamma(\varphi)}{\Gamma(\mu\varphi)\Gamma((1-\mu)\varphi)} y^{\mu\varphi-1} y^{(1-\mu)\varphi-1}, 0 < y < 1, 0 < \mu < 1, \varphi > 0,$$

^۱ نام و ایمیل ارائه دهنده مقاله: کبری قلی زاده، k.gholizadeh@modares.ac.ir

بازپارامتری کرد. در مدل رگرسیونی بتا با ثابت در نظر گرفتن پارامتر دقت φ ، تابعی از میانگین متغیر پاسخ با استفاده از متغیرهای تبیینی موردنظر به صورت $g(\mu) = \mathbf{Z}'\beta$ مدل می‌شود، که در آن $\mathbf{Z}$ مشاهدات متغیرهای تبیینی، β بردار اثرات ثابت و $g(\cdot)$ تابعی یکنواخت از μ است که تابع پیوند نامیده می‌شود. گاهی ممکن است پارامتر دقت بر حسب مشاهدات ثابت نباشد، اسمیتسون و ورکالین (۲۰۰۶) مدل‌بندی پارامتر φ را هم‌زمان با مدل‌بندی پارامتر میانگین پیشنهاد کردند. برای تضمین مثبت بودن واریانس، برای پارامتر دقت تابع پیوند لگاریتم در نظر گرفته می‌شود. در این صورت مدل رگرسیونی بتا به صورت

$$h(\varphi) = \ln(\varphi) = \mathbf{W}'\gamma,$$

خواهد بود، که در آن $\gamma = (\gamma_0, \dots, \gamma_k)$ بردار اثرات ثابت و $\mathbf{W} = (\mathbf{w}_1, \dots, \mathbf{w}_k)$ ماتریس مشاهدات متغیرهای تبیینی با k بعد است.

در بعضی مسائل کاربردی، ممکن است داده‌ها، وابسته فضایی باشند. روش معمول برای لحاظ کردن وابستگی فضایی در مدل‌بندی داده‌های فضایی استفاده از تابع کواریانس فضایی است. در صورت قرار گرفتن موقعیت‌های فضایی مشاهدات بر روی شبکه می‌توان از مدل‌های قدم زدن تصادفی^۱ این همبستگی را در تحلیل داده‌های فضایی لحاظ کرد. در این روش با استفاده از ویژگی مارکوفی و ماتریس همسایگی ساختار کواریانس قابل تعریف است. یک روش برای بیان ویژگی مارکوفی استفاده از گراف است. گراف G مجموعه‌ای از رأس‌هاست که توسط خانواده‌ای از یال‌ها به هم وصل شده‌اند و به صورت زوج مرتب (V, E) نشان داده می‌شود، که در آن V مجموعه‌ای متناهی و ناتهی از رئوس و E یال‌های آن است. مدل‌های قدم زدن تصادفی یک میدان تصادفی مارکوفی گاوسی^۲ (GMRF) است. با توجه به تعریف گراف، بردار تصادفی $\mathbf{X} = (X_1, \dots, X_n)$ یک GMRF تحت گراف $G = (V, E)$ با میانگین μ و ماتریس دقت $\mathbf{Q}_{n \times n} > 0$ است، اگر و تنها اگر تابع چگالی آن به صورت

$$\pi(\mathbf{x}) = (2\pi)^{-\frac{n}{2}} |\mathbf{Q}|^{-\frac{1}{2}} \exp\left\{-\frac{1}{2}((\mathbf{x} - \mu)^T \mathbf{Q}(\mathbf{x} - \mu))\right\},$$

باشد، به گونه‌ای که $Q_{ij} \neq 0$ اگر و تنها اگر $i, j \in E$ (رو و هلد، ۲۰۰۵).

گام‌رمن و چیپدا (۲۰۱۳) با وارد کردن اثرات تصادفی مدل رگرسیون بتا فضایی را مطرح کردند. گاهی در تحلیل بیزی این مدل‌ها توزیع پسینی فرم بسته‌ای ندارند و معمولاً از الگوریتم‌های MCMC استفاده می‌شود. مدل‌بندی هم‌زمان دو پارامتر میانگین و دقت، همچنین وجود همبستگی فضایی در مدل، موجب کاهش کارایی این الگوریتم‌ها و افزایش زمان محاسبات می‌شود. برای حل این مساله می‌توان در تحلیل بیزی مدل رگرسیونی بتا از روش تقریب لاپلاس آشیانی جمع‌بسته^۳ (INLA) استفاده کرد. ارائه تقریب‌های دقیق چگالی پسینی و افزایش سرعت محاسبات از مزیت‌های روش INLA نسبت به الگوریتم‌های MCMC است (رو و همکاران، ۲۰۰۹).

در این مقاله برای تحلیل بیزی مدل رگرسیونی بتا فضایی از روش INLA استفاده می‌شود. برای این منظور در بخش ۲ روش INLA در تحلیل بیزی مدل رگرسیونی بتا ارائه می‌شود. سپس در بخش ۳ با اجرای شبیه‌سازی لزوم مدل‌بندی هم‌زمان پارامتر میانگین و دقت بررسی می‌شود. در بخش ۴ نحوه استفاده مدل رگرسیونی بتا در تحلیل داده‌های فضایی با مثالی کاربردی تشریح می‌شود.

¹ Random Walk

² Gaussian Markov Random Field

³ Integrated Nested Laplace Approximation

۲ برآورد مدل رگرسیونی بتا با INLA

فرض کنید $i = 1, \dots, n$ ، $Y_i \sim \text{Beta}(\mu_i \varphi_i, \varphi_i(1 - \varphi_i))$ بطوری که هر دو پارامتر میانگین μ و دقت φ برای همه مشاهدات ثابت نیستند و به صورت

$$\text{logit}(\mu_i) = \mathbf{z}_i' \boldsymbol{\beta} + \nu_i, \quad \ln(\varphi_i) = \mathbf{z}_i' \boldsymbol{\gamma}, \quad i = 1, \dots, n$$

مدل بندی شوند، که در آن بردار ضرایب رگرسیونی پارامتر میانگین و $\boldsymbol{\beta} = (\beta_1, \dots, \beta_m)$ بردار ضرایب رگرسیونی پارامتر دقت، $\boldsymbol{\gamma} = (\gamma_1, \dots, \gamma_k)$ بردار ν_i بردار اثرات تصادفی در مدل میانگین است. در تحلیل بیزی این مدل با روش INLA بردار اثرات ثابت مدل دقت به عنوان ابرپارامتر در نظر گرفته می شود. برای بردارهای $\boldsymbol{\beta}$ و $\boldsymbol{\nu}$ پیشینی گاوسی در نظر گرفته و فرض می شود $\mathbf{x} = \{\boldsymbol{\nu}, \boldsymbol{\beta}\}$ مشخص کننده بردار همه متغیرهای پنهان است. توزیع چگالی میدان تصادفی $\mathbf{x}$ گاوسی با میانگین صفر و ماتریس دقت $Q(\boldsymbol{\theta}_1)$ است که در آن $\boldsymbol{\theta}_1$ بردار ابرپارامتر با بعد m است. با فرض آنکه مولفه های میدان پنهان $\mathbf{x}$ به شرط $\boldsymbol{\theta}_1$ مستقل هستند، میدان تصادفی پنهان $\mathbf{x}$ یک GMRF با ماتریس دقت تنک $Q(\boldsymbol{\theta}_1)$ است.

فرض کنید متغیرهای پاسخ $\{Y_i : i = 1, \dots, n\}$ به شرط $\mathbf{x}$ و $\boldsymbol{\gamma}$ مستقل هستند و توزیع آن با $\pi(\mathbf{y}|\mathbf{x}, \boldsymbol{\gamma})$ نشان داده می شود. برای راحتی $\boldsymbol{\theta} = (\boldsymbol{\gamma}^T, \boldsymbol{\theta}_1^T)$ با بعد $\dim(\boldsymbol{\theta}) = m + k$ در نظر گرفته می شود. در این صورت، چگالی پسینی $\mathbf{x}$ و $\boldsymbol{\theta}$ به صورت

$$\begin{aligned} \pi(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y}) &\propto \pi(\boldsymbol{\theta})\pi(\mathbf{x}|\boldsymbol{\theta}_1) \prod_{i=1}^n \pi(y_i|\mathbf{x}_i, \boldsymbol{\gamma}) \\ &\propto \pi(\boldsymbol{\theta})|Q(\boldsymbol{\theta}_1)|^{\frac{1}{2}} \exp\left\{-\frac{1}{2}\mathbf{x}^T Q(\boldsymbol{\theta}_1)\mathbf{x} + \sum_{i=1}^n \log\{\pi(y_i|\mathbf{x}_i, \boldsymbol{\gamma})\}\right\}, \end{aligned}$$

خواهد بود. در روش محاسباتی INLA، هدف اصلی توزیع پسینی کناری $\pi(x_i|\mathbf{y})$ و $\pi(\theta_j|\mathbf{y})$ است که می توانند به صورت

$$\begin{aligned} \pi(x_i|\mathbf{y}) &= \int \pi(x_i|\mathbf{y}, \boldsymbol{\theta})\pi(\boldsymbol{\theta}|\mathbf{y})d\boldsymbol{\theta}, \\ \pi(\theta_j|\mathbf{y}) &= \int \pi(\boldsymbol{\theta}|\mathbf{y})d\boldsymbol{\theta}_{-j}, \end{aligned}$$

بازنویسی شوند. برای برآورد این پسینی ها به روش INLA مراحل زیر انجام می شود:

(i) محاسبه $\pi(\boldsymbol{\theta}|\mathbf{y})$ ، که از طریق آن می توان همه پسینی های کناری $\pi(\theta_j|\mathbf{y})$ را بدست آورد.

(ii) محاسبه $\pi(x_i|\mathbf{y}, \boldsymbol{\theta})$ ، که برای محاسبه پسینی کناری $\pi(x_i|\mathbf{y})$ مورد نیاز است.

روش INLA با توجه به مفروضات مدل، تقریب های عددی بر اساس روش تقریب لاپلاس برای برآورد پسینی های مورد نظر ارائه می دهد (ترنی و کادان، ۱۹۸۶).

مرحله (i) شامل محاسبه تقریب پسینی توام ابرپارامترها به صورت

$$\begin{aligned}\pi(\theta|\mathbf{y}) &= \frac{\pi(\mathbf{x}, \theta|\mathbf{y})}{\pi(\mathbf{x}|\theta, \mathbf{y})} \\ &= \frac{\pi(\mathbf{y}|\mathbf{x}, \theta)\pi(\mathbf{x}, \theta)}{\pi(\mathbf{y})\pi(\mathbf{x}|\theta, \mathbf{y})} \frac{1}{\pi(\mathbf{x}|\theta, \mathbf{y})} \\ &= \frac{\pi(\mathbf{y}|\mathbf{x}, \theta)\pi(\mathbf{x}|\theta)\pi(\theta)}{\pi(\mathbf{y})\pi(\mathbf{x}|\theta, \mathbf{y})} \frac{1}{\pi(\mathbf{x}|\theta, \mathbf{y})} \\ &\propto \frac{\pi(\mathbf{y}|\mathbf{x}, \theta)\pi(\mathbf{x}|\theta)\pi(\theta)}{\pi(\mathbf{x}|\theta, \mathbf{y})} \\ &\approx \frac{\pi(\mathbf{y}|\mathbf{x}, \theta)\pi(\mathbf{x}|\theta)\pi(\theta)}{\tilde{\pi}(\mathbf{x}|\theta, \mathbf{y})} \Big|_{\mathbf{x}=\mathbf{x}^*(\theta)},\end{aligned}$$

است، که در آن $\tilde{\pi}(\mathbf{x}|\theta, \mathbf{y})$ تقریب گاوسی چگالی شرطی کامل $\mathbf{x}$ ، و $\mathbf{x}^*(\theta)$ مد این چگالی است (رو و هلد، ۲۰۰۵).
در مرحله (ii)، توزیع‌های پسینی شرطی $\pi(x_i|\mathbf{y}, \theta)$ می‌توانند با استفاده از تقریب لاپلاس به صورت

$$\tilde{\pi}_{LA}(x_i|\theta, \mathbf{y}) \propto \frac{\pi(\mathbf{x}, \theta, \mathbf{y})}{\tilde{\pi}_{GG}(\mathbf{x}_{-i}|x_i, \theta, \mathbf{y})} \Big|_{\mathbf{x}_{-i}=\mathbf{x}_{-i}^*(x_i, \theta)},$$

برآورد شوند، که در آن $\tilde{\pi}_{GG}$ تقریب گاوسی $\pi_G(\mathbf{x}_{-i}|x_i, \theta, \mathbf{y})$ و $\mathbf{x}_{-i}^*(x_i, \theta)$ مد آن است. بعد از انجام دو مرحله، چگالی پسینی کناری $\pi(x_i|\mathbf{y})$ به صورت

$$\tilde{\pi}(x_i|\mathbf{y}) = \sum_k \tilde{\pi}(x_i|\theta_k, \mathbf{y}) \tilde{\pi}(\theta_k|\mathbf{y}) \Delta_k,$$

بدست می‌آید، که در آن Δ_k وزن‌های متناظر با نقاط انتگرال‌گیری θ_k است.

در این مقاله برای مقایسه برآزش و پیچیدگی مدل‌ها از دو ملاک انحراف اطلاع^۴ (DIC) و ملاک اطلاع واتانابه آکایک^۵ (WAIC) استفاده می‌شود. فرض کنید چگالی پسینی است، که در آن θ پارامتر مدل است. ملاک DIC به صورت

$$\begin{aligned}\text{DIC} &= -2 \log \pi(y|\hat{\theta}) + 2(\log \pi(y|\hat{\theta}) - E(\log(\pi(y|\theta)))) \\ &= -2 \log \pi(y|\hat{\theta}) + 2p_{\text{DIC}},\end{aligned}$$

تعریف می‌شود، که در آن p_{DIC} تعداد پارامترهای فعال مدل است (اشپیگلتر و همکاران، ۲۰۰۲). ملاک WAIC که پیچیدگی و برآزش مدل را اندازه‌گیری می‌کند به صورت

$$\text{WAIC} = -2 \log \pi(y|\hat{\theta}) + 2p_{\text{WAIC}},$$

محاسبه می‌شود، که در آن $p_{\text{WAIC}} = 2 \sum_{i=1}^n (\log E(\pi(y|\hat{\theta})) - E(\log(\pi(y|\hat{\theta}))))$ است (واتانابه، ۲۰۱۰). برای مقایسه مدل‌ها، از نظر اعتبار پیشگویی y_i می‌توان از ملاک پیشگویی شرطی مولفه‌ها^۶ (CPO) (CPO) به صورت $CPO_i = \pi(y_i|y_{-i})$ استفاده کرد، که در آن y_{-i} همه نمونه بجز y_i است. در اینجا از ملاک

$$\text{LS} = -\frac{1}{n} \sum_{i=1}^n \log CPO_i$$

⁴Deviance Information Criterion

⁵Watanabe-Akaike Information Criterion

⁶Conditional Predictive Ordinates

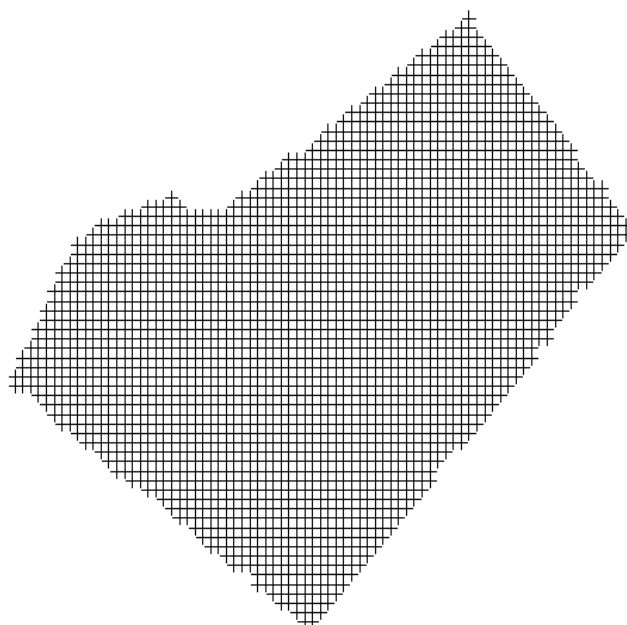
نیز برای مقایسه مدل‌ها استفاده می‌شود، که نمره لگاریتم اعتبارسنجی متقابل^۷ (LS) نامیده می‌شود (نایتینگ و همکاران، ۲۰۰۷).

۳ مطالعه شبیه‌سازی

در این مدل‌بندی پارامتر دقت در مدل‌های رگرسیونی بتا با واریانس‌های مختلف مورد ارزیابی عددی قرار می‌گیرد. مجموعه داده از توزیع بتا $Y_i \sim \text{Beta}(\mu_{st}, \varphi_t)$ ، بر روی شبکه‌ای با 2574 ، $s = 1, \dots$ موقعیت و $t = 1, \dots, 4$ تکرار تولید می‌شود، که در آن پارامتر میانگین و دقت به صورت

$$\text{logit}(\mu_{st}) = i'_t \beta + \beta_x x_1 + f_s(s), \quad \log(\varphi_t) = i'_t \gamma,$$

مدل شده‌اند. قابل ذکر است داده‌ها بر روی شبکه‌ای مشابه موقعیت‌های داده واقعی بخش بعد که نمای کلی آن در شکل ۱ نشان داده شده است، شبیه‌سازی شده‌اند.



شکل ۱: شبکه موقعیت‌ها برای شبیه‌سازی داده‌ها.

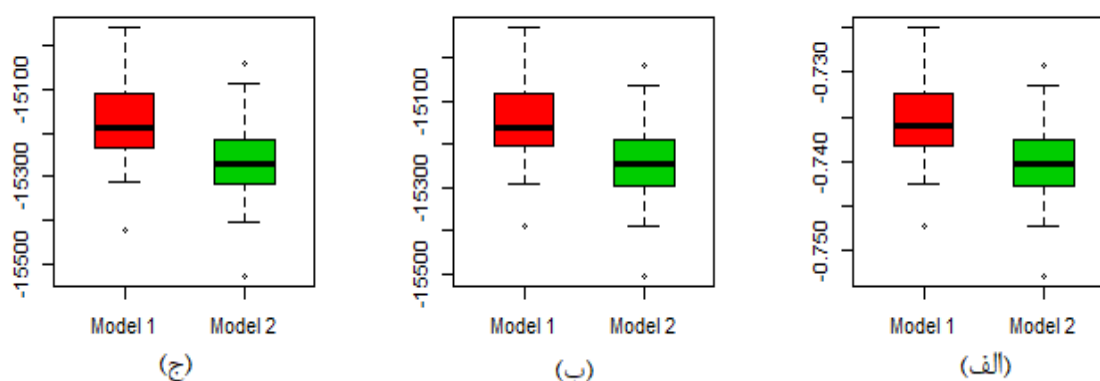
برای محاسبه پارامترهای μ_{st} و φ_t ، مقادیر واقعی $(\gamma_1, \gamma_2, \gamma_3, \gamma_4) = (25, 20, 30, 35)$ در نظر گرفته شده‌اند. اثر تصادفی فضایی $f_s(s)$ از مدل قدم زدن تصادفی مرتبه اول روی شبکه نمونه گرفته شده است. متغیر تبیینی x از توزیع نرمال با میانگین ۰ و واریانس 0.4 تولید شده است. سپس با محاسبه مقادیر μ_{st} و φ_t ، از توزیع $\text{Beta}(\mu_{st}, \varphi_t)$ نمونه Y_i را تولید نموده و این کار ۱۰۰ بار تکرار می‌شود. در هر تکرار دو مدل، اولی با واریانس یکسان در تکرارها (مدل ۱) و دومی با واریانس متفاوت در تکرارها (مدل ۲) برآورد پارامترها با استفاده از INLA محاسبه می‌شود. به عبارت دیگر، در هر دو مدل میانگین یکسان است ولی دقت در مدل ۱ به صورت $\log(\varphi_t) = \gamma$ و در مدل ۲ به صورت $\log(\varphi_t) = i'_t \gamma$ در نظر گرفته شده است. میانگین و انحراف استاندارد برآورد پارامترها در این ۱۰۰ تکرار در جدول ۱ نشان داده شده است. همان طور که در جدول ۱ ملاحظه می‌شود پارامترها

⁷Cross Validated Logarithmic Score

در مدل ۲ در بسیاری موارد با دقت بیشتری نسبت به مدل ۱ برآورد شده‌اند. برای مقایسه بهتر این دو مدل، معیارهای DIC، WAIC و LS در این شبیه‌سازی برای ۱۰۰ تکرار محاسبه شده است. نمودار جعبه‌ای این معیارها در شکل ۲ نمایش داده شده است. همانطور که ملاحظه می‌شود مقادیر سه معیار برای مدل ۱ بیشتر است، یعنی عملکرد مدل ۲ بهتر از مدل ۱ است.

جدول ۱: میانگین پارامترهای برآورد شده مدل‌ها در ۱۰۰ تکرار.

مدل	پارامتر	مقدار واقعی	میانگین برآورد	انحراف استاندارد	چندک	
					۰/۰۲۵	۰/۹۷۵
۱	β_x	۱	۰/۸۷۵	۰/۰۶۸	۰/۷۴۲	۱/۰۰۸
	β_1	۰	-۰/۰۰۱	۰/۰۱۲	-۰/۰۲۵	۰/۰۲۱
	β_2	۱	۰/۹۷۵	۰/۰۱۳	۰/۹۳۲	۰/۹۸۳
	β_3	۱/۵	۱/۳۹۸	۰/۰۱۴	۱/۳۶۱	۱/۴۱۷
	β_4	۲	۱/۸۱۲	۰/۰۱۶	۱/۷۸۱	۱/۸۴۳
۲	β_x	۱	۰/۸۷۹	۰/۰۶۸	۰/۷۴۵	۰/۰۱۳
	β_1	۰	-۰/۰۰۲	۰/۰۱۲	-۰/۰۲۶	۰/۰۲۲
	β_2	۱۱	۰/۹۴۱	۰/۰۱۴	۰/۹۱۴	۰/۹۶۹
	β_3	۱/۵	۱/۴۰۴	۰/۰۱۴	۱/۳۷۷	۱/۴۳۲
	β_4	۲	۱/۸۵۵	۰/۰۱۶	۱/۸۲۴	۱/۸۵۵

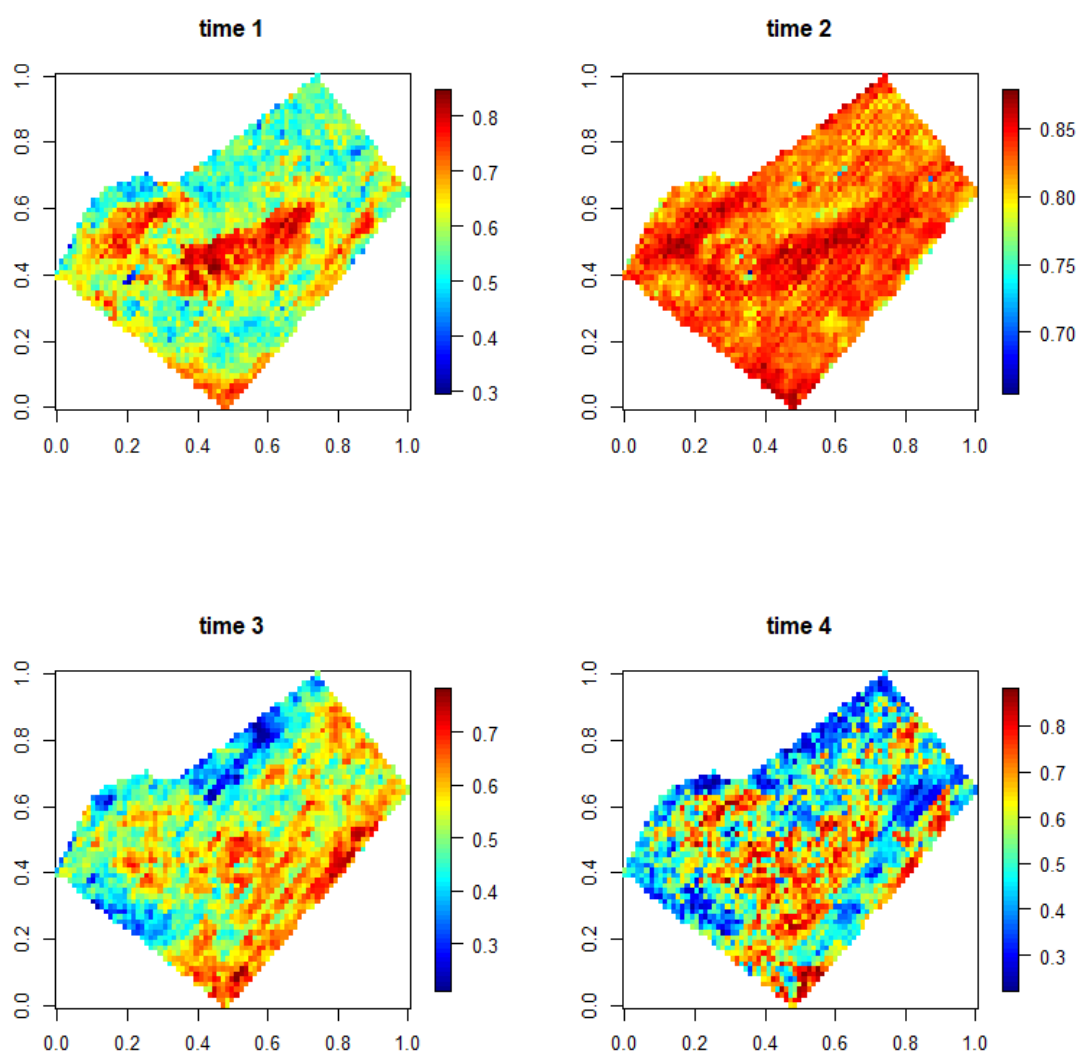


شکل ۲: معیارهای ارزیابی دو مدل الف- LS ب- WAIC ج- DIC.

۴ مدل‌بندی داده‌های رشد گیاه و ویژگی‌های خاک

داده‌های مورد مطالعه در این بخش در مورد استفاده از تکنولوژی حسگرها برای تحلیل مکانی و زمانی روابط خاک و گیاه در اجرای شیوه‌های کارآمد مدیریت کشاورزی است (پولیس و همکاران، ۲۰۱۹). هدف مدل‌بندی روابط بین رشد گیاه

و ویژگی خاک است که توسط دو حسگر در مزرعه یونجه به وسعت هفت هکتار در پالمونته در ایتالیا جنوبی اندازه‌گیری می‌شود. متغیرهای پاسخ و تبیینی به ترتیب شاخص پوشش گیاهی^۸ (NDVI) و مقاومت الکتریکی خاک^۹ (ER) است. موقعیت مکانی شامل مشبکه مربعی به تعداد ۲۵۷۴ موقعیت است. اندازه‌گیری‌های میدانی NDVI چهار مرتبه در فصول مختلف و ER فقط یک بار در سه لایه به عمق ۰/۵، ۱ و ۲ متر صورت گرفته است. شکل ۳ پهنه‌بندی NDVI در چهار زمان را نشان می‌دهد. همان‌طور که ملاحظه می‌شود NDVI دارای بالاترین مقدار و کمترین پراکندگی در مرتبه زمانی دوم است. در شکل ۴ پهنه‌بندی ER در سه عمق نمایش داده شده است. با توجه به این نمودار ER دارای تنوع سیستماتیک قوی در طول سه عمق نیست بنابراین از میانگین آن در این سه لایه در مدل‌بندی استفاده خواهد شد.



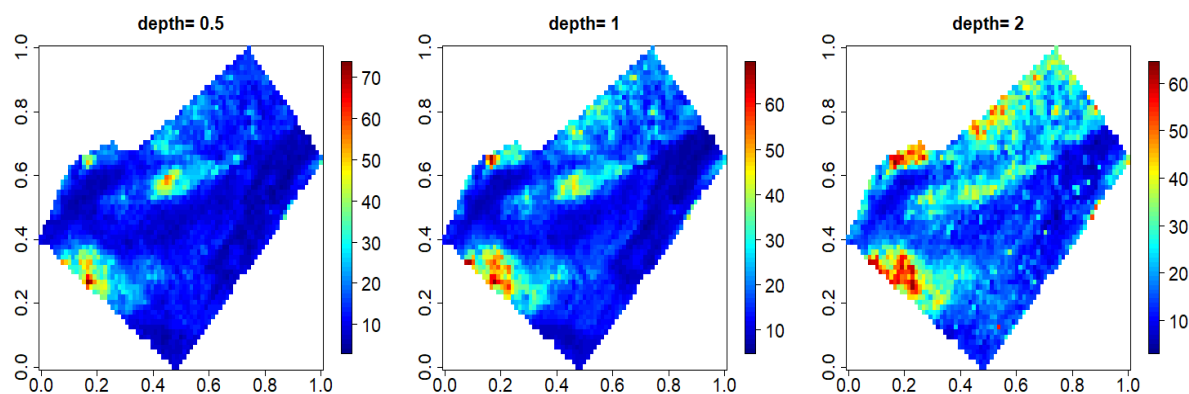
شکل ۳: پهنه‌بندی اندازه‌گیری‌های میدانی در چهار زمان مختلف.

برای بررسی ارتباط بین NDVI و ER مدل دو مدل در نظر گرفته می‌شود، در مدل اول تنها پارامتر میانگین به صورت

$$\text{logit}(\mu_{st}) = \beta_0^\mu + \mathbf{i}_t' \beta_1^\mu + f_1^\mu(ER_s) + f_2^\mu(s), \quad (1.4)$$

⁸The Normalized Density Vegetation Index

⁹The Soil Electrical Resistivity



شکل ۴: پهنه‌بندی میزان مقاومت الکتریکی خاک در سه عمق مختلف.

جدول ۲: اندازه معیارهای ارزیابی در دو مدل.

	LS	WAIC	DIC	
مدل ۱	-۱/۵۵۹۱	-۳۲۸۰۶/۷۱	-۳۳۰۶۳/۴۴	
مدل ۲	-۱/۸۱۸۶	-۳۷۹۵۰/۲۶	-۳۷۸۳۴/۲۲	

مدل می‌شود، که در آن i'_t متغیر تبیینی ظاهری زمان با ۴ سطح، $f_1^{\mu}(\cdot)$ تابع هموار غیرخطی از متغیر ER پنهان و $f_4^{\mu}(\cdot)$ تابع هموار غیرخطی از موقعیت‌ها برای برآورد اثر فضایی است. مدل پیشینی قدم‌زدن تصادفی مرتبه اول برای $f_1^{\mu}(\cdot)$ و مدل پیشینی قدم‌زدن تصادفی دو بعدی برای $f_4^{\mu}(\cdot)$ در نظر گرفته شده‌اند. پارامتر دقت در ۴ تکرار ثابت در نظر گرفته شده است. همانطور که در شکل ۳ ملاحظه می‌شود، واریانس NDVI در چهار تکرار متفاوت است. بنابراین علاوه بر مدل کردن میانگین به صورت (۱.۴) پارامتر دقت نیز به صورت

$$\ln(\varphi_{st}) = i'_t \gamma, \quad (2.4)$$

مدل می‌شود، که در آن i'_t متغیر ظاهری ۴ سطحی زمان است. برای مقایسه این دو مدل معیارهای DIC و WAIC و LS محاسبه و در جدول ۲ نشان داده شده‌اند. همان‌طور که ملاحظه می‌شود با توجه به هر سه معیار عملکرد مدل ۲ از مدل ۱ بهتر است.

بحث و نتیجه‌گیری

در این مقاله مدل رگرسیونی بتا فضایی با مدل‌بندی هم‌زمان پارامتر میانگین و دقت ارائه شد. همان‌طور که در بخش شبیه‌سازی مشاهده شد در صورت ثابت نبودن پارامتر دقت بر حسب متغیر تبیینی نیاز است که این پارامتر علاوه بر پارامتر میانگین مدل شود. در بخش داده واقعی دو مدل برای تحلیل داده‌های پوشش گیاهی در نظر گرفته شد. مدل میانگین هر دو مدل یکسان بود اما پارامتر دقت در ۴ تکرار در مدل اول ثابت و در مدل دوم متفاوت در نظر گرفته شد. در نهایت مشاهده شد که مدل با واریانس متفاوت براساس معیارهای در نظر گرفته شده مدلی بهتر است که این نتیجه با توجه به تحلیل اکتشافی داده‌ها نتیجه منطقی است.

مراجع

- Ferrari S. and Cribari F. (2004), Beta Regression for Modelling Rates and Proportions, *Journal of Applied Statistics*, **31**, 799-815.
- Gamerman D. and Cepeda-Cuervo E. (2013), Generalized Spatial Dispersion Models, *Technical report*, Universidade Federal do Rio de Janeiro.
- Gneiting, T. and Raftery, A. E., (2007), Strictly Proper Scoring Rules, Prediction, and Estimation, *Journal of the American Statistical Association*, **102**, 359–378.
- Rue H., Martino S. and Chopin N., (2009), Approximate Bayesian Inference for Latent Gaussian Models Using Integrated Nested Laplace Approximations (with discussion), *Journal of the Royal Statistical Society*, **71**, 319-392.
- Rue H. and Held I. (2005), *Gaussian Markov Random Field: Theory and Application*. vol 104 of Monographs on Statistics and Applied Probability. Chapman and Hall, London.
- Smithson, M. and Verkuilen, J. (2006), A Better Lemon Squeezer? Maximum-Likelihood Regression with Beta-Distributed Dependent Variables, *Psychological Methods*, **11**, 54-71.
- Spiegelhalter, D. J., Best, N. G., Carlin, B. P., and Van Der Linde, A. (2002), Bayesian Measures of Model Complexity and Fit, *Journal of the Royal Statistical Society: Series B*, **64**, 583– 639.
- Tierney, L. and Kadane, J. (1986), Accurate Approximations for Posterior Moments and Marginal Densities, *Journal of the American Statistical Association*, **81**, 82-86.
- Pollice, A., Lasinio, G., Rossi, R., Amato, M., Kneib, T., and Lang, S. (2019), Bayesian Measurement Error Correction in Structured Additive Distributional Regression with an Application to the Analysis of Sensor Data on Soil-Plant Variability, *Stochastic Environmental Research and Risk Assessment*, **33**, 747-763.
- Watanabe, S. (2010), Asymptotic Equivalence of Bayes Cross Validation and Widely Applicable Information Criterion in Singular Learning Theory, *Journal of Machine Learning Research*, **11**, 3571–3594.



ارتقای الگوریتم چندهدفه ازدحام ذرات برای نمونه‌گیری بهینه فضایی

الهه لطفیان^۱، محسن محمدزاده

گروه آمار، دانشگاه تربیت مدرس

چکیده: تاکنون مطالعات زیادی در رابطه با بهینه‌سازی تک‌هدفه نمونه‌گیری فضایی انجام شده است، ولی در بسیاری از مسائل واقعی، بیش از یک هدف برای نمونه‌گیری فضایی حائز اهمیت است. با توجه به خلا مطالعاتی موجود، استفاده از روش‌های بهینه‌سازی چندهدفه به عنوان راهکاری مناسب برای حل مسائل بهینه‌سازی نمونه‌گیری فضایی با چند هدف پیشنهاد شده است. در این مقاله، یک مسئله بهینه‌سازی دوهدفه برای نمونه‌برداری از رسوبات ساحلی رودخانه میوز جهت ارزیابی میزان آلودگی آن به فلزات سنگین در نظر گرفته شده است. توابع هدف مورد مطالعه، واریانس میانگین نمونه و مسافت طی شده بین موقعیت‌های نمونه هستند. برای این منظور، الگوریتم چندهدفه ازدحام ذرات برای نمونه‌گیری فضایی، ارتقاء داده شده و همبستگی فضایی داده‌ها در الگوریتم لحاظ شده است. همچنین، حساسیت مقادیر توابع هدف نسبت به اندازه نمونه و عملکرد الگوریتم بر اساس معیارهای ارزیابی مورد بررسی قرار گرفته است. نتایج حاصل بیانگر عملکرد مناسب الگوریتم چندهدفه ازدحام ذرات ارتقا یافته در بهینه‌سازی این مسئله از نمونه‌گیری فضایی است.

واژه‌های کلیدی: نمونه‌گیری فضایی، بهینه‌سازی چندهدفه، الگوریتم چندهدفه ازدحام ذرات.

کد موضوعی ریاضی (۲۰۱۰): 90C29, 62H11

۱ مقدمه

یکی از ارکان اساسی جمع‌آوری اطلاعات برای دستیابی به نتایج آماری معتبر، تعیین طرح نمونه بهینه است. در مواجهه با داده‌هایی که همبستگی آن‌ها ناشی از موقعیت قرارگیری در یک فضای معین است و به اصطلاح همبسته فضایی هستند، نمی‌توان از رویکردهای مرسوم نمونه‌گیری استفاده کرد (محمدزاده، ۱۳۹۸). در این حالت باید همبستگی فضایی داده‌ها در طراحی نمونه‌گیری لحاظ شود. طیف گسترده‌ای از مطالعات کاربردی از جمله بررسی خاک، محیط زیست، هواشناسی و غیره با داده‌های فضایی سروکار دارند که نیازمند استفاده از روش‌های نمونه‌گیری فضایی هستند. انتخاب موقعیت‌های

^۱ نام و ایمیل ارائه دهنده مقاله: الهه لطفیان، e.lotfian@modares.ac.ir

نمونه‌گیری در ناحیه مورد مطالعه به منظور بهینه‌سازی برحسب معیارهای ارزیابی، نمونه‌گیری فضایی بهینه نامیده می‌شود. با توجه به ماهیت مسئله، معیارهای مختلفی برای ارزیابی و اخذ نمونه‌های فضایی بهینه تعریف شده است. واریانس پیشگو و آنتروپی پیشگو از جمله معیارهای مورد توجه هستند. همچنین معیارهای دیگری مانند هزینه و زمان نیز بر نحوه نمونه‌گیری اثر گذارند. از آنجا که این معیارها لزوماً هم‌راستا نیستند، برای بهینه‌سازی هم‌زمان برحسب چند معیار، نمی‌توان از روش‌های مرسوم بهینه‌سازی تک‌هدفه استفاده کرد. استفاده از روش‌های بهینه‌سازی چندهدفه می‌تواند راهکاری مناسب برای حل این مشکل باشد. اگر چه مطالعات بسیاری در زمینه بهینه‌سازی تک‌هدفه نمونه‌گیری فضایی صورت گرفته، اما توجه زیادی به بهینه‌سازی چندهدفه نشده است. **مک براتی و همکاران (۱۹۸۱)** تحت فرض تغییرنگار معلوم، با هدف مینیمم کردن واریانس کریگیدن طرح نمونه بهینه را ارائه دادند. **مک براتی و وبستر (۱۹۸۳)** به جای مینیمم کردن واریانس کریگیدن، مینیمم کردن واریانس کریگیدن بلوکی را مورد بررسی قرار دادند. **ساکس و شیلر (۱۹۸۸)** چندین الگوریتم مبتنی بر نوردیدن برای بهینه‌سازی طرح‌های نمونه فضایی ارائه کردند. **گرونکین و استین (۱۹۹۸)** روش نوردیدن شبیه‌سازی شده فضایی ^۱ (SSA) را برای نمونه‌گیری فضایی معرفی کردند. **لارک (۲۰۰۲)** از این روش با هدف برآورد بهینه تغییرنگار برای دستیابی به طرح نمونه استفاده کرد. **مولر و همکاران (۲۰۰۴)** با استفاده از شبیه‌سازی زنجیر مارکوفی ناهمگن و الگوریتم SSA به طرح نمونه بهینه دست یافتند. **فیوننس و همکاران (۲۰۰۷)** با استفاده از یک معیار بهینگی مبتنی بر ماکسیمم آنتروپی و الگوریتم نوردیدن شبیه‌سازی شده فضایی طرح نمونه بهینه را به دست آوردند. **بیداریخت و همکاران (۱۳۹۳)** یک طرح نمونه‌گیری فضایی بهینه برای پیشگویی میزان بارندگی در استان خوزستان با استفاده از معیار آنتروپی تعیین کردند. **گودز و همکاران (۲۰۱۱)** با استفاده از الگوریتم SSA و الگوریتم ژنتیک هیبریدی با هدف مینیمم کردن میانگین واریانس پیشگو در داده‌های زمین آماری به طرح نمونه بهینه دست یافتند. **مارچانت و همکاران (۲۰۱۳)** از روش SSA برای طراحی نمونه‌گیری چندمرحله‌ای خاک‌های آلوده شهری استفاده کردند. **وانگ و همکاران (۲۰۱۴)** این روش را برای برآورد ناحیه‌ای محصول ناخالص اولیه بکار بردند. همچنین، این روش توسط **گومزکاردناس و همکاران (۲۰۱۵)** برای نمونه‌گیری فضایی بهینه در طول آزمایش غیرمخرب سازه‌های بتنی مورد استفاده قرار گرفت. در این مطالعات مسئله مورد بررسی به صورت تک‌هدفه بیان شده است. اما در بسیاری از مسائل کاربردی، بیش از یک هدف برای نمونه‌گیری اهمیت دارد و گاهی این اهداف در تقابل با یکدیگر هستند. یکی از رویکردهای مرسوم برای حل این مسائل، استفاده از روش‌های اسکالرسازی^۲ است که توابع هدف را به یک تابع هدف تبدیل می‌کنند و به این ترتیب، مسئله با استفاده از روش‌های بهینه‌سازی تک‌هدفه قابل حل است (**گرکو و همکاران، ۲۰۰۵**). گرچه این روش‌ها به علت سادگی محاسباتی، به طور گسترده در ادبیات موضوع استفاده شده‌اند، اما در بسیاری از مواقع قادر نیستند توزیع مناسبی از جواب‌های بهینه را مشخص کنند و تنها زمانی که مرز پارتو محدب است می‌توانند به خوبی مرز را پوشش دهند (**ارگوت و همکاران، ۲۰۰۵**). به منظور اجتناب از این مشکلات، الگوریتم‌های بهینه‌سازی چندهدفه مورد توجه قرار گرفته‌اند. **لارک (۲۰۱۶)** الگوریتم بهینه‌سازی چندهدفه مبتنی بر نوردیدن شبیه‌سازی شده^۳ (AMOSA) را برای نمونه‌گیری فضایی توسعه داد. **لطفیان و محمدزاده (۱۳۹۷)** این الگوریتم را برای دست یافتن به جواب‌های بهینه استوار ارتقا دادند.

هدف اصلی این مقاله، ارتقای الگوریتم بهینه‌سازی چندهدفه ازدحام ذرات^۴ (MOPSO) در نمونه‌گیری فضایی و بررسی عملکرد آن است. برای این منظور، در بخش ۲ برخی مفاهیم اولیه در زمینه بهینه‌سازی چندهدفه، الگوریتم MOPSO و معیارهای ارزیابی عملکرد الگوریتم ارائه می‌شود. در بخش ۳، مطالعه موردی و نتایج حاصل از بکارگیری الگوریتم چندهدفه ازدحام ذرات ارتقا یافته برای نمونه‌گیری فضایی تشریح شده است. در نهایت، به بحث و نتیجه‌گیری پرداخته

^۱ Spatial Simulated Annealing

^۲ Scalarization

^۳ Archived Multiobjective Simulated Annealing

^۴ Multi-objective Particle Swarm Optimization

می‌شود.

۲ بهینه‌سازی چندهدفه

یک مسئله بهینه‌سازی چندهدفه در حالت کلی به صورت

$$\begin{aligned} \min f(x) &= (f_1(x), \dots, f_m(x)), \\ g_j(x) &\leq \bullet, \quad j \in J = \{1, \dots, l\}, \\ x &\in X \subseteq \mathbb{R}^n, \end{aligned} \quad (1.2)$$

بیان می‌شود که در آن $f_i, g_j : \mathbb{R}^n \rightarrow \mathbb{R}, (i \in I := \{1, \dots, m\}, j \in J)$ و $m > 1$ هستند. مجموعه ناتهی

$$S := \{x \in \mathbb{R}^n : x \in X, g_j(x) \leq \bullet, j \in J\},$$

مجموعه شدنی و $\mathbb{R}^m$ فضای هدف نامیده می‌شوند. به دلیل وجود چند تابع هدف، بعد فضای تصویر بیشتر از یک بوده و اعمال یک ترتیب کلی استاندارد روی آن ناممکن است (ارگوت و همکاران، ۲۰۰۵). با توجه به این چالش، تعریف کردن جواب بهینه به فرم معمول برای مسائل بهینه‌سازی چندهدفه میسر نیست. بدین منظور، مفهوم جواب بهینه پارتو^۵ مورد توجه قرار می‌گیرد.

تعریف ۱.۲. در مسئله مینیم‌سازی (۱.۲)، جواب شدنی $\hat{x} \in X$ بهینه پارتو (نامغلوب) نامیده می‌شود هرگاه جواب $x \in X$ وجود نداشته باشد که برای $k \in \{1, \dots, m\}$ و برای i ای، $f_i(x) < f_i(\hat{x})$ و $f_k(x) \leq f_k(\hat{x})$.

روش‌های گوناگونی برای یافتن جواب‌های پارتو وجود دارد. اکثر این روش‌ها، توابع هدف را تبدیل به یک تابع هدف می‌کنند که در اصطلاح به روش‌های اسکالرسازی معروفند. این روش‌ها به علت سادگی در حل، در ادبیات موضوع به‌طور مکرر استفاده می‌شوند، ولی دارای اشکالاتی نیز می‌باشند. به عنوان مثال، اگر چه روش جمع وزنی در بسیاری از مسائل بهینه‌سازی چندهدفه برای تولید مرز پارتو مورد استفاده قرار می‌گیرد، ولی تنها زمانی که مرز پارتو محدب است می‌تواند از تمامی بخش‌های مرز نماینده داشته باشد. علاوه بر این، حتی برای مرز پارتو محدب نیز مجموعه وزن‌های مساوی، مجموعه پارتویی با توزیع مساوی تولید نمی‌کنند و قادر نیستند توزیع مناسبی از جواب‌های پارتو را مشخص کنند. بنابراین، استفاده از این روش‌ها ممکن است با مشکلاتی روبرو شود و نتایج مبتنی بر آن‌ها از دقت کافی برخوردار نباشند. با توجه به اشکالات ذکر شده در رابطه با روش‌های اسکالرسازی و همچنین محدودیت‌های موجود در طراحی مدل به صورت تک‌هدفه، استفاده از الگوریتم‌های بهینه‌سازی چندهدفه مرسوم شده است. الگوریتم MOPSO از جمله الگوریتم‌های پرکاربرد در حل بسیاری از مسائل بهینه‌سازی چندهدفه است (ژو و همکاران، ۲۰۱۱). در ادامه، روش‌های انتخاب بولتزمن^۶ و چرخ رولت^۷ معرفی می‌شوند که در روند الگوریتم مورد استفاده قرار می‌گیرند (کوئلو و همکاران، ۲۰۰۴).

۱.۲ روش‌های انتخاب مبتنی بر شایستگی

۱.۱.۲ روش بولتزمن

فرض کنید شرط‌های زیر برای احتمال انتخاب p_i برقرار باشند:

^۵Pareto

^۶Boltzmann

^۷Roulette Wheel

$$0 \leq p_i \leq 1 \quad ۱.$$

$$\sum_{i=1}^n p_i = 1 \quad ۲.$$

۳. اگر $n_i \leq n_j$ آنگاه $p_i \geq p_j$ (اصل شایسته‌سالاری).

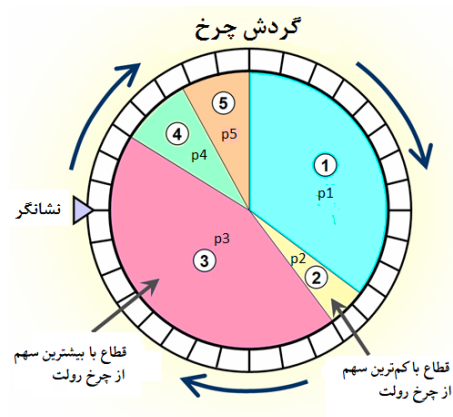
از اصل شایسته‌سالاری می‌توان نتیجه گرفت که ناحیه‌های با تعداد اعضای بیشتر احتمال کمتری برای انتخاب دارند. به عبارتی، در این روش به هر ناحیه، احتمال p_i به صورت

$$p_i = \frac{e^{-\beta n_i}}{\sum_j e^{-\beta n_j}}$$

نسبت داده می‌شود، که در آن n_i تعداد اعضای خانه i ام و β پارامتر فشار انتخاب^۸ است.

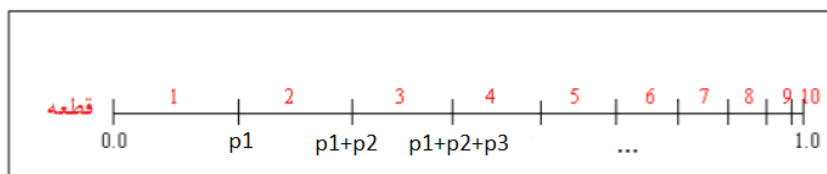
۲.۱.۲ روش چرخ رولت

در چرخ رولت، برای هر عضو، یک قطاع از چرخ اختصاص داده می‌شود. قطاع‌ها در اندازه‌های مختلف و متناسب با میزان اهمیت عضو متناظر هستند. بزرگ‌ترین قطاع به مهم‌ترین عضو و کوچک‌ترین قطاع به کم اهمیت‌ترین عضو تخصیص می‌یابد. چرخ چرخانده می‌شود و عضو مرتبط با قطاع برگزیده انتخاب می‌شود. همان‌طور که در شکل ۱ ملاحظه می‌شود، احتمال توقف نشانگر در قطاع i ام برابر p_i است (پاروز ۲۰۱۴). برای سادگی بیشتر، در شکل ۲، دایره به صورت پاره خط



شکل ۱: چرخ رولت

نشان داده شده است. با تولید یک عدد تصادفی r در بازه $[0, 1]$ ، یک قطاع انتخاب می‌شود. به‌طوری‌که، اگر $r \geq 0$ و



شکل ۲: پاره‌خط متناظر با چرخ رولت

⁸Selection Pressure Parameter

$r \leq p_1$ ، قطاع اول، اگر $r > p_1 + p_2$ و $r \leq p_1$ ، قطاع دوم و روند به همین شکل برای انتخاب سایر قطاع‌ها ادامه پیدا می‌کند.

۲.۲ الگوریتم بهینه‌سازی چندهدفه ازدحام ذرات

الگوریتم MOPSO ابتدا توسط **کوئلو و همکاران (۲۰۰۴)** معرفی شد. در این الگوریتم، یک آرشیو^۹ خارجی برای ذخیره پاسخ‌های نامغلوبی که تاکنون تولید شده‌اند، به کار می‌رود. آرشیو، شامل دو بخش مهم کنترل‌کننده آرشیو و شبکه‌بندی است و مهم‌ترین هدف آن نگهداری از جواب‌های نامغلوب یافت شده طی فرآیند جست‌وجو می‌باشد. کنترل‌کننده آرشیو تعیین می‌کند که آیا پاسخ جدیدی باید به آرشیو اضافه شود یا نه. فرآیند تصمیم‌گیری به این صورت است که پاسخ‌های نامغلوبی که در هر تکرار الگوریتم به دست می‌آیند، با اعضای آرشیو که در ابتدا تهی است مقایسه می‌شوند. اگر آرشیو خالی باشد، در این صورت پاسخ‌های فعلی قابل قبول‌اند. اگر پاسخ جدید توسط عضوی از آرشیو مغلوب شود، این پاسخ حذف و اگر هیچ یک از اعضای آرشیو، پاسخ جدید را مغلوب نکند، این پاسخ در آرشیو ذخیره می‌شود. سرانجام، اگر جمعیت آرشیو به ظرفیت ماکسیمم خود برسد، برای ایجاد مرزهای پارتو با توزیع مناسب، از روش شبکه‌بندی تطبیقی استفاده می‌شود. در این رویکرد، فضای توابع هدف به چند ناحیه تقسیم می‌شود، اگر عضوی از آرشیو خارج از مرزهای فعلی شبکه قرار گیرد، شبکه باید بار دیگر محاسبه شده و هر عضو آن دوباره موقعیت‌دهی شود. سرعت و موقعیت هر ذره در این الگوریتم، به ترتیب با استفاده از رابطه‌های

$$VEL[i] = W \times VEL[i] + R_1(PBESTS[i] - POP[i]) + R_2(REP[h] - POP[i]),$$

$$POP[i] = POP[i] + VEL[i].$$

محاسبه و به‌روز می‌شوند، که در آن W وزن اینرسی، R_1 و R_2 اعداد تصادفی در بازه $[0, 1]$ و $PBESTS[i]$ بهترین موقعیت ذره i ام هستند. $REP[h]$ مقداری است که از آرشیو به دست آمده و اندیس h به این صورت انتخاب می‌شود که پس از شبکه‌بندی و محاسبه تعداد اعضای خانه‌های جدول، برای هر خانه به کمک روش بولتزمن، احتمالی در نظر گرفته می‌شود، به طوری که خانه‌های با تعداد اعضای بیشتر احتمال کمتری برای انتخاب دارند. پس از تعیین احتمال انتخاب، یک خانه به کمک روش چرخ رولت مشخص شده و در نهایت یکی از اعضای آن به تصادف انتخاب می‌شود. روند کلی الگوریتم به صورت زیر است:

۱. ایجاد جمعیت اولیه (POP).
۲. مقداردهی اولیه به سرعت هر ذره.
۳. ارزیابی هر ذره از جمعیت.
۴. جداکردن اعضای نامغلوب جمعیت و ذخیره آن‌ها در آرشیو.
۵. شبکه‌بندی فضای هدف کشف شده.
۶. انتخاب رهبر توسط هر ذره از میان اعضای آرشیو.
۷. به‌روز کردن بهترین موقعیت هر یک از ذرات.

^۹Archive

۸. اضافه کردن اعضای نامغلوب جمعیت فعلی به آرشیو،

۹. حذف کردن اعضای مغلوب آرشیو،

۱۰. اگر تعداد اعضای آرشیو بیش از ظرفیت تعیین شده باشد، اعضای اضافی نیز حذف می‌شوند (اندازه آرشیو محدود است)،

۱۱. اگر شرایط خاتمه محقق نشده باشد، به مرحله ۵ برمی‌گردیم و در غیر این صورت پایان.

۳.۲ معیارهای ارزیابی عملکرد الگوریتم

برای مقایسه جواب‌های بهینه بدست آمده توسط الگوریتم از لحاظ پوشش و همگرایی مرز پارتو، از معیارهایی مانند میانگین فاصله ایده‌ال^{۱۰} (MID)، متر فاصله^{۱۱} (SM)، متر پراکندگی^{۱۲} (DM) و گسترش جواب‌های نامغلوب^{۱۳} (SNS) استفاده می‌شود (مقصودلو و همکاران، ۲۰۱۶).

میانگین فاصله ایده‌ال (MID): برای اندازه‌گیری فاصله جواب‌های مرزهای پارتو از یک نقطه ایده‌ال، از این معیار استفاده می‌شود. این معیار به صورت

$$MID = \frac{1}{n} \sum_{i=1}^n \left[\left(\frac{f_{\lambda i} - f_{\lambda best}}{f_{\lambda total}^{max} - f_{\lambda total}^{min}} \right)^2 + \left(\frac{f_{\nu i} - f_{\nu best}}{f_{\nu total}^{max} - f_{\nu total}^{min}} \right)^2 \right]^{1/2},$$

تعریف می‌شود، که در آن n تعداد جواب‌های نامغلوب، $f_{\lambda total}^{max}$ و $f_{\lambda total}^{min}$ به ترتیب بزرگ‌ترین و کوچک‌ترین مقادیر توابع هدف اول و دوم در بین تمام جواب‌های نامغلوب، $f_{\nu total}^{min}$ و $f_{\nu total}^{max}$ به ترتیب کوچک‌ترین و بزرگ‌ترین مقادیر توابع هدف اول و دوم در بین تمام جواب‌های نامغلوب و $f_{\lambda best}$ و $f_{\nu best}$ به ترتیب بهترین مقادیر توابع هدف اول و دوم در بین تمام جواب‌های نامغلوب هستند.

متر فاصله (SM): برای محاسبه یکنواختی جواب‌های نامغلوب از این معیار استفاده می‌شود که با استفاده از فرمول

$$SM = \frac{\sum_{i=1}^n |\bar{d} - d_i|}{(n-1)\bar{d}},$$

محاسبه می‌شود، که در آن n تعداد جواب‌های نامغلوب به دست آمده توسط الگوریتم، d_i فاصله اقلیدسی بین دو جواب نامغلوب متوالی مرز پارتو و $\bar{d}$ میانگین این فاصله‌ها است. الگوریتم با مقادیر کوچک‌تر MID و SM عملکرد بهتری دارد. متر پراکندگی (DM): این معیار برای محاسبه پراکندگی جواب‌های پارتو تعریف شده است که از رابطه

$$DM = \left[\sum_{i=1}^m (\min f_i - \max f_i)^2 \right]^{1/2},$$

محاسبه می‌شود، که در آن $\min f_i$ و $\max f_i$ به ترتیب کم‌ترین و بیش‌ترین مقدار هر تابع هدف در بین تمام جواب‌های نامغلوب هستند.

گسترش جواب‌های نامغلوب (SNS): از این معیار برای محاسبه پراکندگی و تنوع جواب‌های پارتو استفاده می‌شود. این

¹⁰Mean Ideal Distance

¹¹Spacing Metric

¹²Diversification Metric

¹³Spread of Non-dominance Solution

معیار از رابطه

$$SNS = \left[\frac{\sum_{i=1}^n (MID - C_i)^2}{n-1} \right]^{1/2},$$

به دست می آید، که در آن $C_i = [f_1^2 + f_2^2]^{1/2}$. الگوریتم با مقادیر بزرگتر DM و SNS دارای عملکرد بهتری است.

۳ مطالعه موردی

آلودگی زیست بومها به فلزات سنگین، یکی از مهم ترین مسائل محیط زیست است که زندگی گیاهان، جانوران و به ویژه انسانها را تهدید می کند. در حال حاضر، بسیاری از رودخانه ها در معرض آلودگی فلزات سنگین ناشی از منابع طبیعی همچون هوازدگی و فرسایش سنگ و همچنین منابع انسانی از جمله پساب های شهری، صنعتی و کشاورزی قرار دارند. رسوبات به عنوان مخزن آلاینده ها، شاخص مهمی برای تعیین آلودگی آب محسوب می شوند. این مطالعه، درصدد تعیین مکان های مناسب نمونه برداری در طول سواحل رودخانه میوز، به منظور ارزیابی هر چه دقیق تر غلظت فلزات سنگین موجود در رسوبات ساحلی این رودخانه است. در این پژوهش، منطقه مورد مطالعه برای داده های میوز در نظر گرفته شده است که اولین بار توسط **بارو و مک دونالد (۱۹۹۸)** ارائه شد. منطقه مورد نظر شامل یک دشت سیلابی در طول رودخانه میوز واقع در هلند به طول جغرافیایی ۵۲ درجه و ۰ دقیقه و ۵۲۸۸۵۲/۱۳ ثانیه شمالی و به عرض جغرافیایی ۴ درجه و ۴۶ دقیقه و ۵۲۲۸۰۰/۲۲ ثانیه شرقی است. در این مقاله، به منظور تعیین موقعیت های مناسب نمونه برداری، واریانس میانگین نمونه و مسافت طی شده برای نمونه برداری به عنوان توابع هدف مسئله فرض می شوند. برآورد مدل محور واریانس میانگین نمونه، به عنوان تابع هدف اول به صورت

$$\sigma_m^2 = \frac{2}{n} \sum_{i=1}^n \bar{\gamma}(\mathbf{x}_i, B) - \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n \gamma(\mathbf{x}_i - \mathbf{x}_j) - \bar{\gamma}(B, B),$$

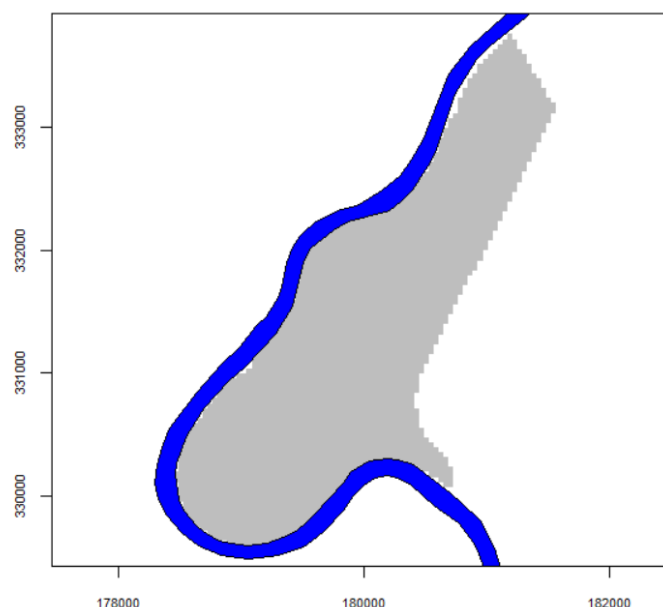
محاسبه می شود، که در آن بردار موقعیت i امین نمونه، $\gamma(\mathbf{h})$ نیم تغییرنگار و B دامنه نمونه گیری است (**ویستر و الیور ۱۹۹۰**). همچنین، $\bar{\gamma}(\mathbf{x}_i, B)$ و $\bar{\gamma}(B, B)$ به ترتیب، به صورت

$$\begin{aligned} \bar{\gamma}(\mathbf{x}_i, B) &= \int_{\mathbf{x}_k \in B} \gamma(\mathbf{x}_i - \mathbf{x}_k) d\mathbf{x}_k, \\ \bar{\gamma}(B, B) &= \int_{\mathbf{x}_k \in B} \int_{\mathbf{x}_\ell \in B} \gamma(\mathbf{x}_k - \mathbf{x}_\ell) d\mathbf{x}_\ell d\mathbf{x}_k, \end{aligned}$$

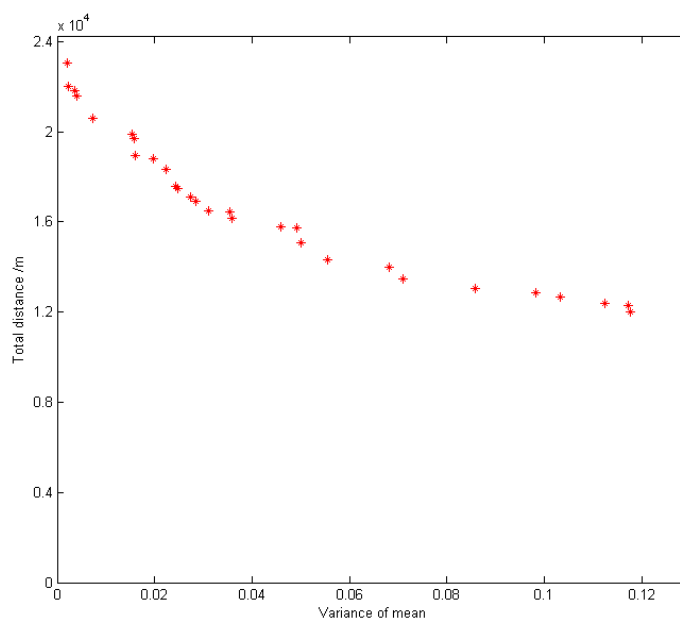
محاسبه می شوند. تابع هدف دوم، مسافت طی شده بین موقعیت های نمونه است که برای مینیمم سازی آن باید کوتاه ترین مسیر بین نقاط نمونه را به دست آورد. برای این منظور، از روش فروشنده دورگرد^{۱۴} (TSP) استفاده می شود (**گوتین و پون، ۲۰۰۶**). تابع هدف اول توسط طرح نمونه ای مینیمم می شود که بهترین پوشش فضایی را در سطح منطقه مورد مطالعه ایجاد کند، در حالی که چنین طرح نمونه ای تابع هدف دوم را ماکسیمم می کند و این بیانگر تقابل بین دو تابع هدف است. در این مطالعه، برای یافتن طرح نمونه بهینه، الگوریتم چندهدفه ازدحام ذرات برای داده های فضایی ارتقا یافته است. با توجه به ماهیت فضایی داده ها، ساختار همبستگی فضایی در تابع هدف واریانس میانگین نمونه لحاظ می شود. برای داده های میوز، متغیر مورد نمونه گیری دارای همبستگی فضایی با تغییرنگار کروی با دامنه ۸۹۷ متر، اثر قطعه ای ۰/۵۱ و ازاره ۰/۵۹۱ است و ۱۵۵ موقعیت برای نمونه برداری از خاک سطحی منطقه به منظور بررسی غلظت فلزات سنگین انتخاب می شود (**پسما**،

¹⁴ Travelling Salesman Problem

۲۰۱۹). در ادامه، نتایج حاصل از اجرای الگوریتم برای مسئله مطرح شده با ۲۰۰۰ تکرار ارائه می‌شود. برنامه‌ها روی یک کامپیوتر با واحد پردازش مرکزی (CPU)، ۳۰/۲ GHz، با پنج هسته، حافظه دستیابی تصادفی (RAM)، ۸/۰۰ GB و نرم‌افزار MATLAB b۲۰۱۳R اجرا شده‌اند.



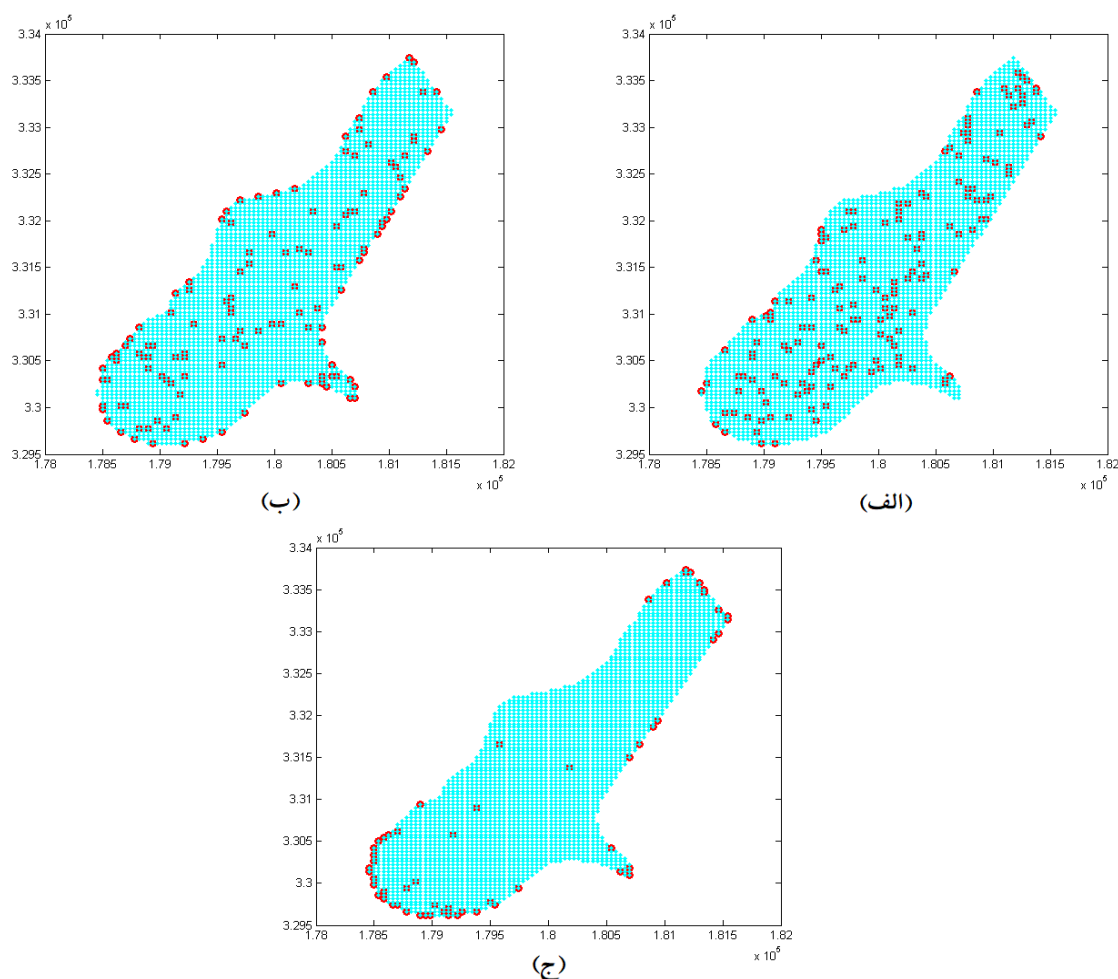
شکل ۳: منطقه مورد مطالعه در طول رودخانه میوز



شکل ۴: مرز پارتو تقریبی به دست آمده توسط الگوریتم

شکل‌های ۳ و ۴ به ترتیب منطقه مورد مطالعه و مرز پارتو تقریبی بدست آمده توسط الگوریتم را نشان می‌دهند. همان‌طور که از شکل ۴ ملاحظه می‌شود، مرز پارتو به دست آمده توسط الگوریتم، پراکندگی نسبتاً خوبی دارد. شکل ۵ نیز

برخی از موقعیت‌های نمونه‌برداری به‌دست آمده توسط الگوریتم را نشان می‌دهد. در جدول ۱ مقادیر توابع هدف در برخی



شکل ۵: برخی از موقعیت‌های نمونه‌برداری به‌دست آمده توسط الگوریتم

از طرح‌های نمونه محاسبه شده است. همان‌طور که ملاحظه می‌شود، با افزایش مقدار واریانس، مسافت طی شده کاهش می‌یابد و برعکس، که این خود گویای تقابل بین دو تابع هدف است.

جدول ۱: مقادیر توابع هدف در برخی از طرح‌های نمونه (با اندازه نمونه ۱۵۵)

طرح نمونه	(الف)	(ب)	(ج)
واریانس میانگین نمونه	۰۰۲۲/۰	۰۱۱۹/۰	۱۰۹۸/۰
مسافت طی شده (متر)	۲۳۵۳۳	۲۱۴۳۹	۱۳۲۷۲

در جدول ۲ مقادیر هر یک از معیارهای ارزیابی عملکرد الگوریتم ارائه شده است. ملاحظه می‌شود معیارهای SM و MID مقادیر کم و معیارهای DM و SNS مقادیر بالایی دارند که نشان‌دهنده پوشش و پراکندگی مناسب مرز پارتو تقریبی به‌دست آمده و عملکرد مناسب الگوریتم است.

جدول ۲: مقادیر معیارهای ارزیابی عملکرد الگوریتم (با اندازه نمونه ۱۵۵)

الگوریتم	SM	MID	DM	SNS
MOPSO	۸۴۶۴/۰	۶۵۴۹/۰	۱۰۲۶۱	۱۸۹۱۷

جدول ۳: ماکسیمم و مینیمم مقدار توابع هدف را برای اندازه نمونه‌های مختلف نشان می‌دهد. همان‌طور که انتظار می‌رود، با افزایش اندازه نمونه، واریانس میانگین نمونه کم‌تر و مسافت طی شده بیشتر می‌شود.

جدول ۳: ماکسیمم و مینیمم مقدار توابع هدف برای اندازه نمونه‌های متفاوت

اندازه نمونه	واریانس میانگین نمونه		مسافت طی شده (متر)	
	مینیمم	ماکسیمم	مینیمم	ماکسیمم
۵۰	۰/۰۰۴۸	۰/۱۲۰۹	۱۰۵۷۸	۱۴۹۲۱
۱۰۰	۰/۰۰۳۳	۰/۱۱۰۲	۱۱۷۰۴	۱۹۳۹۴
۱۵۵	۰/۰۰۲۲	۰/۱۰۹۸	۱۳۲۷۲	۲۳۵۳۳

بحث و نتیجه‌گیری

یک مسئله بهینه‌سازی دوهدفه برای نمونه‌برداری از رسوبات ساحلی رودخانه میوز جهت ارزیابی میزان آلودگی آن به فلزات سنگین در نظر گرفته شد. واریانس میانگین نمونه و مسافت طی شده بین موقعیت‌های نمونه به عنوان توابع هدف مسئله فرض شدند. در این راستا، الگوریتم MOPSO برای یافتن طرح نمونه بهینه فضایی ارتقا داده شد. برای این منظور ساختار همبستگی فضایی داده‌ها در الگوریتم لحاظ شد. نتایج حاصل نشان‌دهنده عملکرد مناسب الگوریتم MOPSO ارتقا یافته در این مسئله از نمونه‌گیری فضایی است.

مراجع

- بیداربخت، ح.، محمدزاده، م.، محمدی، الف. و مرید، س. (۱۳۹۳)، طراحی نمونه‌گیری فضایی بهینه با ملاک آنتروپی، دومین کارگاه آموزشی اندازه‌های اطلاعات و کاربردهای آن، ۸-۹ بهمن ۱۳۹۳، مشهد، ایران.
- لطفیان، الف و محمدزاده، م.، (۱۳۹۷)، نمونه‌گیری فضایی با استفاده از بهینه‌سازی چندهدفه استوار، چهاردهمین کنفرانس آمار ایران، شاهرود، ایران، ۵۲۷-۵۳۴.
- محمدزاده، م.، (۱۳۹۸)، آمار فضایی و کاربردهای آن، چاپ سوم، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران.

Burrough, P. A., McDonnell, R. A. (1998), *Principles of Geographical Information Systems*, 2nd Edition. Oxford University Press.

Coello, C. A. C., Pulido, G. T., and Lechuga, M. S. (2004), Handling Multiple Objectives with Particle Swarm Optimization, *IEEE Transactions on evolutionary computation*, **8**, 256-279.

Ehrgott, M., Figueira, J., and Greco, S. (Eds.). (2005), *Multiple Criteria Decision Analysis: State of*

the Art Surveys, Springer.

Fuentes, M., Chaudhuri, A., and Holland, D. M. (2007), Bayesian Entropy for Spatial Sampling Design of Environmental Data, *Journal Environmental Ecological Statistics*, **14**, 323-340.

Gomez-Cardenas, C., Sbartai, Z. M., Balayssac, J. P., Garnier, V., and Breysse, D. (2015), New Optimization Algorithm for Optimal Spatial Sampling During Non-Destructive Testing of Concrete Structures, *Engineering Structures*, **88**, 92-99.

Greco, S., Ehrgott, M., and Figueira, J.R. (Eds). (2005), *Multiple Criteria Decision Analysis*, State of the Art Surveys, Springer, Boston.

Groenigen, J. W., and Stein, A. (1998), Constrained Optimization of Spatial Sampling Using Continuous Simulated Annealing, *Journal of Environmental Quality*, **27**, 1078-1086.

Guedes, L. P. C., Riberio JR, P. J., Piedade, S. M. D., and UribeOpazo, M. A. (2011), Optimization of Spatial Sample Configurations using Hybrid Genetic Algorithm and Simulated Annealing, *Chilean Journal of Statistics*, **2**, 39-50.

Gutin G., and Punnen A. P. (2006), *The Traveling Salesman Problem and its Variations*, Springer Science and Business Media.

Lark, R. M. (2002), Optimized Spatial Sampling of Soil for Estimation of the Variogram by Maximum Likelihood, *Geoderma*, **105**, 49-80.

Lark, R. M. (2016), Multi-objective Optimization of Spatial Sampling, *Spatial Statistics*, **18**, 412-430.

Maghsoudlou, H., Afshar-Nadjafi, B., and Niaki, S. T. A. (2016), A Multi-objective Invasive Weeds Optimization Algorithm for Solving Multi-skill Multi-mode Resource Constrained Project Scheduling Problem, *Computers & Chemical Engineering*, **88**, 157-169.

Marchant, B. P., McBratney, A. B., Lark, R. M., and Minasny, B. (2013), Optimized Multi-Phase Sampling for Soil Remediation Surveys, *Spatial Statistics*, **4**, 1-13.

McBratney, A. B., and Webster, R. (1983), Optimal Interpolation and Isarithmic Mapping of Soil Properties, *European Journal of Soil Science*, **34**, 137-162.

McBratney, A. B., Webster, R., and Burgess, T. M. (1981), The Design of Optimal Sampling Schemes for Local Estimation and Mapping of Regionalized Variables-I: Theory and Method, *Computers and Geosciences*, **7**, 331-334.

Muller, P., Sanso, B., and De Iorio, M. (2004), Optimal Bayesian Design by Inhomogeneous Markov Chain Simulation, *Journal of the American Statistical Association*, **99**, 788-798.

- Parvez, R. (2014), Selection in Evolutionary Algorithm. <https://www.slideshare.net/riyadparvez/selection-in-evolutionary-algorithm>. Accessed 10 May 2018.
- Pebesma, E. (2019). The Meuse Data Set: A Brief Tutorial for the gstat R Package.
- Sacks, J., and Schiller, S. (1988), Spatial Designs. In *Statistical Decision Theory and Related Topics IV*, Gupta, S. S., Berger, J. O. (eds), Springer, New York, 385-399.
- Wang, J. H., Ge, Y., Heuvelink, G. B. M., and Zhou, C. H. (2014), Spatial Sampling Design for Estimating Regional GPP with Spatial Heterogeneities, *IEEE Geoscience Remote Sensing Letters*, **11**, 539-543.
- Webster, R., and Oliver, M. A. (1990), *Statistical Methods in Soil and Land Resource Survey*, Oxford University Press, Oxford.
- Zhou, A., Qu, B. Y., Li, H., Zhao, S. Z., Suganthan, P. N., and Zhang, Q. (2011), Multiobjective Evolutionary Algorithms: A Survey of the State of the Art, *Swarm and Evolutionary Computation*, **1**, 32-49.



مدل‌های توام نقشه‌بندی بیماری‌ها

یدالله محرابی^۱، امیر کاوسی^۱، بهزاد مهکی^۲، پریسا پورنظری^{۱۲}

^۱گروه اپیدمیولوژی، دانشکده بهداشت و ایمنی، دانشگاه علوم پزشکی شهید بهشتی

^۲گروه آمار زیستی، دانشکده بهداشت، دانشگاه علوم پزشکی کرمانشاه

چکیده: نقشه بندی بیماری شامل مجموعه ای از آماری برای تهیه نقشه ی برآوردهای میزان های بیماری ها است. بیشتر مطالعات نقشه بندی بیماری ها بر مدل بندی یک بیماری واحد متمرکز بوده است اما بیماری های زیادی وجود دارند که عوامل خطر مشترکی داشته، در عمل همبسته بوده و استفاده از مدل‌هایی که اطلاعات مربوط به چند بیماری مرتبط را ادغام می‌کنند، می‌تواند مفید باشد. در سال های اخیر چندین مدل برای تحلیل توأم چند بیماری معرفی شده و مطالعات متعددی انجام گرفته که در آن از مدل های نقشه‌بندی چندمتغیره ی بیماری ها استفاده شده است. با این وجود تاکنون به صورت جامع، روش های تحلیل توأم نقشه‌بندی بیماری ها مورد بررسی و مقایسه قرار نگرفته‌اند. در مطالعه حاضر مهم ترین مدل های نقشه بندی چندمتغیره ی بیماری ها معرفی گردیده و مبانی تئوری، کاربردها، محاسن و معایب هر یک به اختصار بیان می‌شوند.

واژه‌های کلیدی: نقشه بندی بیماری ها، مدل های توام، مدل های چند متغیره.

کد موضوع بندی ریاضی (۲۰۱۰): 62H30، 62H11.

۱ مقدمه

نقشه بندی بیماری شامل مجموعه ای از روش های آماری است که به تهیه نقشه های دقیق براساس برآوردهایی از میزان های بیماری ها و در نهایت برآورد توزیع جغرافیایی بیماری مورد نظر منجر می شوند. اهداف اصلی نقشه بندی بیماری عبارتند از: ۱. تشریح تغییرات مکانی یا زمانی در میزان های بیماری یا خطر به منظور تدوین فرضیات علت شناسی ۲. ارزیابی خطر و تعیین نواحی دارای خطر بالا که نیازمند اعمال مداخله هستند. ۳. استفاده از آن‌ها به منظور تخصیص بهتر منابع و خدمات ۴. ساخت اطلس بیماری ها (لاوسون و همکاران، ۲۰۰۳، ۲۰۰۰).

^۱ نام و ایمیل ارائه دهنده مقاله: پریسا پورنظری، p.pournazari93@gmail.com

در سالیان اخیر تلاش های فراوانی برای نقشه بندی میزان های بیماری ها صورت گرفته است. یکی از مهم ترین مشکلات در این خصوص، انتخاب معیار مناسب برای نقشه بندی است. نقشه بندی خام میزان های بیماری ها ممکن است کاملاً گمراه کننده باشد. به منظور حل این مشکل، معمولاً از یکی از معیارهای خطرات نسبی در هر ناحیه که با نسبت مرگ و میر یا نسبت بروز استاندارد شده ی اختصاصی ناحیه SMR^1 یا SIR^2 اندازه گیری می شود، استفاده می گردد. گرچه استفاده از این برآوردها در نقشه بندی بیماری ها بسیار رایج است اما نقاط ضعف بسیاری در استفاده از این شاخص ها وجود دارد که برای حل آن می توان از هموارسازی های آماری برای غلبه بر این مشکلات استفاده کرد (لاوسون و همکاران، ۲۰۰۰). تاکنون چندین مدل برای نقشه بندی تک متغیره ی بیماری ها مطرح شده اند که مهم ترین آن عبارتند از:

۱. مدل های هموارسازی که در جهت هموارسازی نوسانات بر پایه توابعی از داده ها در نواحی مورد نظر تلاش می کنند. رایج ترین روش در این کلاس، روش هموارسازی کرنل نادارایا واتسون است (سیمونوف، ۲۰۱۲). ۲. مدل های بیز خطی که مبتنی بر توابع خطی از SMR یا SIR هستند (مارشال، ۱۹۹۱). ۳. مدل های بیز کامل BYM^3 که از مهم ترین مدل های نقشه بندی بیماری ها هستند. در این مدل پارامتر خطر نسبی بیماری به سه مؤلفه روند، تغییرات همبسته مکانی و تغییرات ناهمبسته تجزیه می شود (بساگ و همکاران، ۱۹۹۱). ۴. مدل های بیز تجربی که مشابه روش های بیزی بوده اما توزیع پیشین خطرات نسبی از داده های مشاهده شده برآورد می شود. این روش ها بینابین روش حداکثر درست نمایی و روش های بیز کامل بوده و فرض می کنند که خطرات نسبی از توزیع های خاصی پیروی می کنند که با مجموعه ای از پارامترها مشخص می شوند. این پارامترها را می توان از داده ها با حداکثر نمودن درستنمایی حاشیه ای برآورد کرد (خوشکار و همکاران، ۲۰۱۵). ۶. مدل های آمیخته که کاملاً از مدل های پیشین که در آن ها پراکندگی و تغییرات در خطر با مخلوطی از چند جزء مدل بندی می شود متفاوت است (روهینگ، ۱۹۹۹). ۷. مدل های چندسطحی که به داده هایی برازش می یابند که دارای سطوح خوشه بندی در ساختار خود هستند. در نقشه بندی بیماری ها و کاربردهای جغرافیایی، این سطوح، سطوح متفاوت جغرافیایی در نظر گرفته می شوند (رستاقی و همکاران، ۲۰۱۵).

با وجود پیشرفت های سریع روش های آماری و محاسباتی، گسترش تحلیل های نقشه بندی بیماری ها اغلب بر مدل بندی یک بیماری واحد متمرکز بوده است حال آن که بیماری های زیادی هستند که عوامل خطر مشترکی دارند و اگر الگوهای مشترکی از پراکندگی جغرافیایی بیماری های مرتبط در تحلیل توأم مشخص گردد آنگاه نسبت به حالت تحلیل جداگانه، شواهد متقاعدکننده تری از خوشه بندی واقعی در خطر به دست می آید. اگر داده های چند بیماری یا رویداد بهداشتی که الگوهای جغرافیایی مشترکی دارند در دسترس باشد مدل بندی توأم خطرات می تواند منجر به بهبود برآوردها نسبت به تحلیل جداگانه ی هر رویداد گردد. حتی در صورتی که یک رویداد به عنوان رویداد هدف مدنظر باشد، ورود داده های دیگر رویداد های مرتبط در دسترس در تحلیل مفید است (نورهلد و همکاران، ۲۰۰۱).

در سال های اخیر، توجه به تحلیل توأم چند بیماری که بالقوه مرتبط هستند افزایش یافته، در این راستا، چندین مدل کلاسیک و بیزی، معرفی شده و مطالعات موردی متعددی انجام گرفته است که در آن از کاربردها و مدل های نقشه بندی چندمتغیره ی بیماری ها استفاده شده است با این وجود جز در چند مطالعه اندک و به صورت محدود، روش های متعدد تحلیل توأم نقشه های بیماری ها با هم مقایسه نشده و شرایط کلی آن ها مورد بررسی قرار نگرفته است. با توجه به اهمیت موضوع و کاربرد رو به گسترش نقشه بندی بیماری ها، در این مقاله مهم ترین مدل های نقشه بندی توأم و چندمتغیره ی بیماری ها معرفی گردیده و مبانی تئوری، کاربردها، محاسن و معایب هر یک به اختصار بیان می گردد.

¹Standardized Mortality/Morbidity Ratio

²Standardized Incidence Ratio

³Besag, York and Mollie

۲ مدل‌های نقشه‌بندی چندمتغیره ی بیماری‌ها

۱.۲ مدل‌های رگرسیون اکولوژیک

این مدل‌ها توسط برناردینلی و همکاران معرفی شده‌اند که در آن میزان‌های یک بیماری ثانویه به عنوان جانشین عوامل خطر در مدل رگرسیونی وارد می‌شود. با استفاده از این رویکرد، یک چارچوب نقشه‌بندی چندمتغیره ی بیزی بیماری‌ها برای بررسی‌های اکولوژیک و مکانی الگوهای میزان‌های بیماری و ارتباطات بین بیماری‌ها و مرگ و میر با عوامل خطر اجتماعی، اقتصادی و محیطی به دست می‌آید. در این روش، ابتدا مدل BYM برای داشتن عوامل خطر در سطوح مکانی x_i بسط داده می‌شود.

$$\log(\theta_i) = \alpha + u_i + \nu_i + x_i\beta$$

که در آن، θ معرف خطر نسبی، u_i و ν_i به ترتیب بیانگر اثرات ناهمگنی همبسته مکانی و ناهمگنی ناهمبسته مکانی در ناحیه ی i و β بیانگر اثر عوامل خطر است. هدف از ورود اطلاعات مواجهه x_i ، تشریح پراکندگی مکانی است. سپس رویکرد رگرسیون اکولوژیک با در نظر گرفتن خطای نمونه‌گیری و خطای اندازه‌گیری عامل خطر در مدل، تعدیل می‌شود بدین ترتیب که اگر z_i نشان دهنده ی عامل خطر واقعی اما مشاهده نشده باشد و از آن به جای x_i استفاده می‌کنیم لگاریتم خطر نسبی عبارت خواهد بود از:

$$\log(\theta_{i1}) = \alpha + u_i + \nu_i + z_i\gamma$$

z_i را به عنوان جانشین مواجهه به بیماری ارتباط می‌دهیم. این کار را می‌توان با فرض کردن مدل خطای اندازه‌گیری $\log(\theta_{i2})|z_i \sim N(z_i, \omega)$ انجام داد که $\log(\theta_{i2})$ لگاریتم خطر نسبی بیماری جانشین و ω دقت خطای اندازه‌گیری است. این مدل با فرض $O_{i2}|\log(\theta_{i2}) \sim \text{poisson}(E_{i2}\theta_{i2})$ تکمیل می‌شود که O_{i2} و E_{i2} به ترتیب تعداد موارد مشاهده و مورد انتظار بیماری می‌باشند (حداد و همکاران، ۲۰۱۵).

۲.۲ مدل‌های چندسطحی

لیلاند و همکاران، تحلیل مکانی توأم دو بیماری را براساس مدل چندسطحی ارائه نموده‌اند که در آن اجزاء نامعلوم واریانس از طریق روش‌های درست‌نمایی برآورد می‌شوند. در این مدل برآورد میزان‌های رویدادها در نواحی کوچک از طریق مدل بندی مکانی داده‌ها و افزودن اطلاعات بیشتر به هر ناحیه بهبود می‌یابد. استفاده از مدل‌های مکانی برای پیش‌بینی همزمان بیش از یک رویداد مد نظر قرار می‌گیرد. این کار با در نظر گرفتن مدل مکانی به شکل چندسطحی و به دنبال آن دربرگرفتن ساختار چندمتغیره ی داده‌ها انجام می‌شود. برآوردها نیز با استفاده از حداقل مربعات تعمیم یافته ی مکرر به دست می‌آیند (لانگفورد و. همکاران، ۱۹۹۹؛ لیند و همکاران، ۲۰۰۰).

۳.۲ مدل‌های نرمال دو متغیره (بیزی و غیربیزی)

چندین مدل با استفاده از توزیع‌های نرمال برای نقشه‌بندی توأم دو بیماری معرفی شده‌اند که برخی از آن‌ها بیزی و برخی دیگر غیربیزی هستند. مولی و دسوزا (مولایی، ۱۹۹۰) رویکردی برای تحلیل توأم تعداد موارد ناحیه ای در نواحی، مبتنی بر پیشین نرمال دو متغیره برای پارامترهای لگاریتم خطر نسبی متناظر در هر ناحیه اتخاذ کردند اما این مدل از همبستگی مکانی ممکن خطرات نسبی در نواحی، که توسط لانگفورد و همکاران (لانگفورد و. همکاران، ۱۹۹۹) در نظر گرفته شده است، چشم‌پوشی می‌کند. دیگر محققین توزیع اتورگرسیو شرطی CAR^۴ را برای دو بیماری بسط داده‌اند یا مستقیماً از توزیع

^۴Conditional Autoregressive

های چندمتغیره ی اتورگرسیو شرطی استفاده کرده اند. کیم و همکاران نیز جایگزینی بیزی را برای تحلیل مکانی توأم بیماری ها پیشنهاد نموده اند که محدود به مدل بندی دو بیماری است (کیم و همکاران، ۲۰۰۱).

۴.۲ مدل های میدان های تصادفی مارکف نرمال چندمتغیره

گلفند و وناتسو مدل بندی توأم با استفاده از میدان های تصادفی مارکف نرمال چندمتغیره را توصیه کرده اند. مدل های میدان های تصادفی نرمال با جایگزینی توزیع شرطی نرمال چندمتغیره به جای توزیع نرمال تک متغیره برای بردار لگاریتم خطرات نسبی در ناحیه i به مدل های میدان های تصادفی نرمال چندمتغیره بسط داده شده و مدل BYM چندمتغیره ساخته شده است (جین و همکاران، ۲۰۰۵؛ گلفند و همکاران، ۲۰۰۳).

۵.۲ مدل های جزء مشترک SC

نور-هلد و بست مدل جزء مشترک را برای تحلیل توأم مکانی دو یا چند بیماری معرفی کردند. از این مدل برای ارزیابی همبستگی مکانی بین دو بیماری، استفاده می شود. ایده کلیدی این مدل جداکردن و تقسیم فضای پایه ی خطر به جزء مشترک دو بیماری و جزء اختصاصی هر کدام از بیماری ها است. جزء مشترک به عنوان جانشینی برای عوامل خطر مشاهده نشده که ساختار مکانی را نشان می دهند و برای هر دو بیماری مشترکند، به کار می رود. جزء اختصاصی هر بیماری نیز تنها معرف و جانشین عوامل خطر پراکنده در سطح مکانی است که مختص همان بیماری است. همچنین وزن نسبی هر جزء مشترک برای بیماری های مرتبط در مدل وارد می شود. مدل به دست آمده قابلیت تعیین بزرگی پراکندگی را از طریق الگوهای جغرافیایی در فضای بیماری ها فراهم می آورد. تاکنون چندین شکل از مدل SC^۵ معرفی شده است که در ادامه مهم ترین آن ها معرفی می گردد (نورهلد و همکاران، ۲۰۰۱).

۱.۵.۲ مدل SC مقارن برای دو بیماری

به منظور مدل بندی توأم دو بیماری در ناحیه ی i ($i = 1, \dots, n$) با موارد مشاهده شده O_{ik} و تعداد موارد مورد انتظار E_{ik} ، مدل

$$O_{ik} \sim \text{Poisson}(E_{ik}\theta_{ik}) \quad k = 1, 2$$

به کار می رود و تجزیه ی زیر برای لگاریتم خطرات نسبی در نظر گرفته می شود:

$$\log(\theta_{i1}) = \alpha_1 + \lambda_1\delta + \phi_{i1}$$

$$\log(\theta_{i2}) = \alpha_2 + \lambda_2\delta + \phi_{i2}$$

که در آن θ_{i1} و θ_{i2} به ترتیب معرف خطر نسبی دو بیماری در ناحیه ی i ، α_k معرف عرض از مبدأ اختصاصی بیماری k و λ_i معرف عامل خطر مشترک دو بیماری و ϕ_{i1} و ϕ_{i2} نیز به ترتیب عوامل خطر اختصاصی بیماری نخست و دوم در ناحیه ی i هستند. δ نیز به منظور فراهم آوردن شیب خط متفاوت ارتباط عامل مشترک با بیماری k به مدل وارد شده است (نورهلد و همکاران، ۲۰۰۱).

⁵Shared Component

۲.۵.۲ مدل SC نامتقارن برای دو بیماری

در این مدل یک جزء مشترک مرتبط با هر دو بیماری مرتبط λ_i و یک جزء اختصاصی ϕ_{i1} و مرتبط با تنها یکی از آن‌ها وجود دارد.

$$\log(\theta_{i1}) = \alpha_1 + \lambda_1 \delta + \phi_{i1}$$

$$\log(\theta_{i2}) = \alpha_2 + \lambda_1 / \delta$$

این فرمول بندی تا اندازه ای از فرمول بندی متقارن برای مدل SC متفاوت است. مدل متقارن دارای ویژگی جالب توجه تقارن است اما معمولاً عدم وجود فرض پیشین قوی درباره ی پارامترهای دقت مشکلاتی را در خصوص قابلیت شناسایی عوامل خطر ایجاد می کند. لگاریتم خطر نسبی مشترک، اشتراکات دو بیماری در هر ناحیه را اندازه گیری می کند (هلد و همکاران، ۲۰۰۵).

۳.۵.۲ مدل SC چندمتغیره

تعمیم مدل SC به بیش از دو بیماری سراسر و آسان است. چندین عامل خطر می توانند برای چند بیماری یا تنها برای یک بیماری به اشتراک گذاشته شوند. فرض می شود لگاریتم خطر نسبی برای ناحیه i و بیماری k دارای میانگین برابر با جمع عرض از مبدأ و جمع وزنی عوامل خطر مربوطه است. برای هر عامل خطر تصادفی مشترک یا اختصاصی با توزیع نرمال فرض می شود که بردار $[\log(\delta_{l,1}), \dots, \log(\delta_{l,n_l})]'$ دارای توزیع نرمال چندمتغیره با میانگین صفر و ماتریس واریانس کواریانس Σ است و n_l بیانگر تعداد بیماری های مرتبط با عامل خطر است (هلد و همکاران، ۲۰۰۵؛ دنینگ و همکاران، ۲۰۰۸). مهکی و همکاران (۲۰۱۱) تحلیل مکانی ۷ نوع سرطان با ۵ عامل خطر مشترک را در ۳۰ استان ایران با استفاده از این مدل به شکل زیر انجام داده اند:

$$\log(\theta_{i1}) = \alpha_1 + \lambda_{1i} \delta_{1,1} + \lambda_{2i} \delta_{2,1} + \lambda_{3i} \delta_{3,1} + \lambda_{4i} \delta_{4,1} + \varepsilon_{i1}$$

$$\log(\theta_{i2}) = \alpha_2 + \lambda_{1i} \delta_{1,2} + \lambda_{3i} \delta_{3,2} + \lambda_{4i} \delta_{4,2} + \varepsilon_{i2}$$

$$\log(\theta_{i3}) = \alpha_3 + \lambda_{1i} \delta_{1,3} + \lambda_{4i} \delta_{4,3} + \varepsilon_{i3}$$

$$\log(\theta_{i4}) = \alpha_4 + \lambda_{2i} \delta_{2,2} + \lambda_{4i} \delta_{4,4} + \lambda_{5i} \delta_{5,1} + \varepsilon_{i4}$$

$$\log(\theta_{i5}) = \alpha_5 + \lambda_{1i} \delta_{1,4} + \lambda_{4i} \delta_{4,5} + \varepsilon_{i5}$$

$$\log(\theta_{i6}) = \alpha_6 + \lambda_{2i} \delta_{2,3} + \lambda_{4i} \delta_{4,6} + \varepsilon_{i6}$$

$$\log(\theta_{i7}) = \alpha_7 + \lambda_{2i} \delta_{2,4} + \lambda_{4i} \delta_{4,7} + \lambda_{5i} \delta_{5,2} + \varepsilon_{i7}$$

که در آن $\theta_{i1}, \dots, \theta_{i7}$ به ترتیب معرف خطرات نسبی سرطان های مری، معده، مثانه، روده بزرگ، ریه، پروستات و پستان است. $\lambda_{1i}, \dots, \lambda_{5i}$ نیز پنج جزء مشترک به ترتیب جانشین عوامل خطر مصرف سیگار، اضافه وزن یا چاقی، معرف مصرف کم میوه و سبزیجات، شاخص توسعه انسانی و فعالیت فیزیکی کم هستند. پارامترهای مقیاس نامعلوم δ برای هر جزء، امکان فراهم آوردن شیب خطر متفاوت برای بیماری ها مرتبط با آن جزء را فراهم می آورند. ε_{ik} نیز اثرات ناهمگنی اختصاصی بیماری ها بوده که بیش پراکندگی پوآسون تشریح شده در مدل را در نظر می گیرد (مهکی و همکاران، ۲۰۱۱).

۶.۲ مدل مرگ و میر نسبی PM

این مدل هنگامی استفاده می شود که تنها تعداد مرگ ناشی از بیماری ها در دسترس بوده و تعداد جمعیت در معرض خطر نامشخص باشد. در این حالت می توان توزیع تعداد مرگ ها برای بیماری مورد نظر نسبت به کل مرگ ها را به دست آورد. به منظور فرمول بندی این مدل، ابتدا فرض می شود $O_{ik} \sim \text{poisson}(\lambda_{ik})$ که اندیس $i = 1, \dots, n$ معرف نواحی مورد مطالعه و اندیس $k = 1, 2$ بیانگر دو بیماری مورد نظر می باشند. با این فرض که $\log(\lambda_{ki}) = \alpha_{0k} + \alpha_{1k}x_i$ به گونه ای که مدل فاقد اثرات تصادفی است و x_i ها نمایانگر عوامل خطر در سطوح ناحیه ای مرتبط با دو بیماری از طریق خطرات نسبی مکانی $\exp(\alpha_{11}) \exp(\alpha_{12})$ هستند. اگر مجموع تعداد موارد دو بیماری در ناحیه i را با $M_i = O_{i1} + O_{i2}$ تعریف کنیم آنگاه $O_{i1} | M_i = m_i \sim \text{Bin}(m_i, \gamma_i)$ که:

$$\gamma_i = \frac{\exp\{(\alpha_{01} - \alpha_{02}) + (\alpha_{11} - \alpha_{12})x_i\}}{1 + \exp\{(\alpha_{01} - \alpha_{02}) + (\alpha_{11} - \alpha_{12})x_i\}}$$

بنابراین می توان داده ها را با مدل لجستیک به فرم $\text{logit}(\gamma_i) = \beta_0 + \beta_1 x_i$ مدل بندی نمود که پارامتر β_1 را می توان به عنوان تفاوت $\alpha_{11} - \alpha_{12}$ تفسیر نمود. بدین ترتیب استنباط می تواند براساس تفاوت در لگاریتم خطرات نسبی بدون داشتن جامعه ی در معرض خطر انجام شود. همچنین توجه شود که هنگامی که بیماری دوم مرتبط با عامل خطر x نیست داریم $\beta_1 = \alpha_{11}$ (دابی و همکاران، ۲۰۰۵).

۷.۲ مدل لجستیک چندگانه PL

این مدل به عنوان بسط مدل PM برای زمانی که بیش از دو بیماری مورد تحلیل قرار می گیرد معرفی شده است که در آن یک بیماری به عنوان گروه پایه و مرجع تعریف شده و برای هر پیشگو، فرمول و دستور مدل SC به کار برده می شود. فرض می شود که $O_{ij} = (O_{ij1}, \dots, O_{ijk}, \dots, O_{ijK})'$ از توزیع $\text{multi}(m_{ij}, \pi_{ij} = (\pi_{ij1}, \dots, \pi_{ijk}, \dots, \pi_{ijK}))'$ پیروی می کند که $m_{ij} = \sum_{k=1}^K O_{ijk}$ و $\sum_{k=1}^K \pi_{ijk} = 1$. مدل PL را به عنوان بسط بیش از دو گروه مدل PM در نظر می گیریم که هر گروه با احتمالی به شکل $\pi_{ijk} = \phi_{ijk} / \sum_{r=1}^K \phi_{ijr}$ مدل بندی می شود و هر لگاریتم شانس، به عرض از مبدأ اختصاصی بیماری (α_k^* معرف تفاوت کلی بین بیماری k و بیماری مرجع K)، عبارت ساختاریافته ی زمان با سن و بیماری (α_{jk}^* معرف تفاوت کلی بین بیماری k و بیماری مرجع K)، اثرات ساختاریافته و ساختاریافته ی مکانی (u_{ik}^* و ν_{ik}^* به ترتیب معرف تفاوت بین عبارات ساختاریافته و ساختاریافته ی فضایی بین بیماری k و بیماری مرجع) تجزیه می شود:

$$\log(\phi_{ijk}) = \alpha_k^* + \alpha_{jk}^* + u_{ik}^* + \nu_{ik}^*$$

در نهایت، از مدل SC به منظور تجزیه عبارات ساختاریافته ی مکانی u_{ik}^* به اثرات مشترک و اختصاصی بیماری استفاده می شود.

$$u_{i2}^* = ua_i \delta_2 + up_{i2} \quad u_{i1}^* = ua_i \delta_1 + up_{i1}$$

که در آن ua_i جزء ساختاریافته ی مکانی مشترک و up_{i1} و up_{i2} اجزاء ساختاریافته ی مکانی اختصاصی بیماری هستند که هدف نهایی تمرکز بر برآورد اثرات این اجزاء به عنوان متغیرهای پنهان بیانگر عوامل خطر اختصاصی بیماری است (درسی، ۲۰۰۷؛ نصرزادانی و همکاران، ۲۰۱۸)

۸.۲ مدل‌های مکانی-زمانی

بسیاری از کاربردهای نقشه‌بندی بیماری‌ها مانند تعیین پراکندگی مکانی خطر بیماری‌ها یا بررسی خوشه‌بندی، معمولاً محدود به یک دوره‌ی زمانی هستند اما در پژوهش‌های پزشکی و بهداشتی، بسیاری از داده‌ها در چندین دوره‌ی زمانی و به مدت چند سال گردآوری می‌شوند. در این صورت می‌توان تحلیل نقشه‌های بیماری‌ها که شامل بعد زمانی است را در نظر گرفت. در سالیان اخیر، پیشرفت قابل توجهی در مدل‌های آماری و تکنیک‌های تحلیل داده‌های مکانی-زمانی رخ داده و مدل‌های مکانی به مدل‌های مختلف مکانی-زمانی میزان‌های بیماری‌ها بسط داده شده است (ریچاردسون و همکاران، ۲۰۰۶؛ احمدی پناه و همکاران، ۲۰۱۸؛ مهکی و همکاران، ۲۰۱۸؛ راعی و همکاران، ۲۰۱۸).

همانند تحلیل‌های مکانی، در تحلیل‌های مکانی-زمانی نیز مدل‌های توأم و چندمتغیره نسبت به مدل‌های تک متغیره برتری‌ها و محاسن متعددی مانند بهبود دقت برآوردها الگوی بیماری‌ها و توانایی تعیین الگوهای مشترک بین بیماری‌ها و در نتیجه تقویت نتایج مربوطه دارند. به همین دلیل بسط مدل‌های مکانی-زمانی تک متغیره به مدل‌های مکانی-زمانی چندمتغیره مورد توجه قرار گرفته است. تاکنون مطالعات اندکی ایده‌ی نقشه‌بندی مکانی-زمانی بیماری‌ها را با تحلیل توأم دو یا بیش از دو بیماری ترکیب کرده و از مدل‌های بیزی سلسله‌مراتبی، بیز کامل، رویکرد تحلیل عاملی، مدل‌های مکانی-زمانی SC و .. استفاده نموده‌اند (احمدی پناه و همکاران، ۲۰۱۸؛ مهکی و همکاران، ۲۰۱۸؛ راعی و همکاران، ۲۰۱۸؛ تی زالا و همکاران، ۲۰۰۳).

۱.۸.۲ مدل بیز کامل مکانی-زمانی چندمتغیره

در این مدل O_{ijk} تعداد موارد بروز بیماری k در ناحیه‌ی i و دوره‌ی زمانی j و θ_{ijk} برآورد خطر نسبی (SMR) در نظر گرفته می‌شود، فرض می‌شود O_{ijk} ‌ها از توزیع پواسون با میانگین $E_{ijk}\theta_{ijk}$ پیروی می‌کنند که E_{ijk} ‌ها بیانگر مقادیر مورد انتظار بیماری‌ها هستند. مدل لگ خطی برای میزان‌های بروز بیماری‌ها به کار گرفته شده و با وارد نمودن اثرات پراکندگی مکانی و زمانی، این مدل تغییرپذیری در θ_{ijk} به دلیل نواحی، سال‌ها و بیماری‌ها را در نظر می‌گیرد. روند مکانی پنهان برای معرفی ارتباط بین نواحی همسایه و روند زمانی پنهان برای مشخص نمودن همبستگی در طول زمان انتخاب می‌شود. ارتباط مکانی-زمانی مورد نظر بر اساس عوامل خطر مشخص می‌گردد. بدین ترتیب مدل عبارت است از:

$$O_{ijk} \sim \text{Poisson}(E_{ijk}\theta_{ijk})$$

$$\log(\theta_{ijk}) = \alpha_k + \delta_k\lambda_i + \psi_k t_j + \varepsilon_{ijk}$$

که در آن K معرف بیماری‌ها و α_k عرض از مبدأ اختصاصی بیماری k است. λ_i نشان دهنده‌ی اثر تصادفی مکانی در ناحیه‌ی i و δ_k معرف پارامتر مقیاس مکانی برای بیماری k است. t_j بیانگر اثر تصادفی زمانی در دوره‌ی زمانی j و ψ_k معرف پارامتر مقیاس زمانی برای بیماری k است. ε_{ijk} پراکندگی در نظر گرفته نشده در مدل است (اوسن و همکاران، ۲۰۰۸).

۹.۲ مدل‌های مکانی-زمانی SC دو متغیره

می‌توان بعد زمانی را از طریق اجزاء مشترک و اختصاصی زمانی به مدل SC وارد نمود. در ساده‌ترین شکل، مدل مکانی-زمانی فاقد عبارات اثر متقابل بین زمان و مکان و ناهمگنی است:

$$\mu_{ij1} = \lambda_i\delta + t_j\psi, \quad \mu_{ij2} = \frac{\lambda_i}{\delta} + \frac{t_j}{\psi} + \phi_i + \gamma_j$$

که در آن λ_i معرف الگوی مکانی مشترک، ϕ_i معرف تفاوت الگوی مکانی مشترک دو بیماری، t_j معرف روند زمانی مشترک و γ_j معرف تفاوت روند زمانی دو بیماری است. ضرائب δ و ψ معرف اثر نسبی سهم عبارات مشترک در خطر بیماری اول در مقایسه با بیماری دوم است. می توان اثرات متقابل مکانی- زمانی مشترک بین دو بیماری و ناهمگنی های اختصاصی بیماری ها را به مدل اضافه نمود (احمدی پناه و همکاران، ۲۰۱۸؛ مهکی و همکاران، ۲۰۱۸؛ راعی و همکاران، ۲۰۱۸).

۱۰.۲ مدل تحلیل عاملی مکانی- زمانی چند متغیره

رویکرد تحلیل عاملی برای در نظر گرفتن تحلیل عاملی پراکندگی مکانی- زمانی در میزان های اختصاصی نواحی چند بیماری در طول زمان بسط داده شده است. در این رویکرد چارچوب مدل سازی سلسله مراتبی بیزی اتخاذ گردیده و فرمول بندی های پیشین متفاوت برای عوامل مکانی و زمانی پنهان منعکس کننده ی الگوی مشترک خطرات در نظر گرفته می شود. چارچوبی پیشنهاد می شود که به آسانی می تواند هر تعداد بیماری مورد نظر را وارد آن نمود ضمن آن که همبستگی زمانی نیز در نظر گرفته می شود. این کار با فرمول بندی مدل در یک چارچوب مدل بندی سلسله مراتبی انجام می گیرد که در آن در گام نخست از درست نمایی پواسون برای تعداد مشاهده شده ی بیماری ها استفاده می شود و در گام دوم عوامل پنهان معرفی می گردند (تی زالا و همکاران، ۲۰۰۳).

۳ مقایسه ی مدل ها

۱.۳ مقایسه مدل های چند متغیره و تک متغیره

تحلیل توأم دو یا چند بیماری دارای محاسن متعددی نسبت به تحلیل تک متغیره است. این محاسن شامل بهبود ثبات برآوردها برای نواحی جداگانه، توانایی تشخیص الگوهای مشترک و اختصاصی بیماری ها و خطرات که در تحلیل جداگانه کمتر آشکار است، مشخص کردن خوشه های توأم مربوط به عوامل خطر مشترک بیماری ها، توانایی تعیین وابستگی مکانی بین چند بیماری و ارتباط آن ها با عوامل خطر و افزایش دقت برآورد خطرات بیماری ها است. به علاوه انتظار می رود مدل بندی توأم تغییرپذیری در برآورد خطرات را کاهش دهد. کاهش تغییرپذیری در بیماری های دارای شیوع و بروز بالاتر، کمتر است اما در بیماری های نادر، میزان کاهش چشمگیر است. همچنین مدل بندی توأم بهبودی قابل ملاحظه ای را بر اساس معیار DIC^۶، نسبت به مدل های تک متغیره نتیجه می دهد. این بهبودی به دلیل بهبود برازش مدل با تعداد کمتر پارامترهای مؤثر مور نیاز در مدل بندی توأم است (هلد و همکاران، ۲۰۰۵؛ دنینگ و همکاران، ۲۰۰۸؛ دابنی و همکاران، ۲۰۰۵).

۲.۳ مقایسه مدل های توأم و رگرسیون اکولوژیک

رویکرد رگرسیون اکولوژیک در کنار مدل های چند سطحی، اولین گام ها در نقشه بندی معرفی تحلیل مکانی توأم دو بیماری بوده است. در ابتدا پیشنهاد می شد که از میزان های بروز دیگر بیماری ها به عنوان اندازه های مواجهه ی جانشین در رگرسیون استفاده شود با این حال مدل بندی توأم که به طور همزمان تغییرات مکانی در خطر دو یا چند بیماری مرتبط با هم را مدل بندی می کند، رویکردی ساده تر، پرتوان تر و طبیعی تر برای تشخیص الگوهای جغرافیایی در خطر به نظر می رسد زیرا بیماری ها به عنوان متغیرهای پاسخ وارد مطالعه شده و قابل ارجاع و انتساب به عوامل خطر پنهان می باشند (هلد و همکاران، ۲۰۰۵).

^۶Deviance Information Criterion

۳.۳ مقایسه مدل‌های چندمتغیره ی بیزی و غیر بیزی

میزان‌های به دست آمده از مدل‌های چندمتغیره ی بیزی دقیق‌تر بوده در حالی که در روش‌های غیربیزی، فواصل اطمینان بسیار بزرگ است. همچنین اجرای روش زنجیر مارکوف مونت کارلو ساده‌تر بوده و برآوردهای قابل اعتمادتری را ارائه می‌دهند (اسونسا و همکاران، ۲۰۰۴؛ ماندا و همکاران، ۲۰۰۷).

۴.۳ مقایسه مدل‌های SC و نرمال چند متغیره MVN

بر اساس شاخص DIC مدل SC بهتر از مدل MVN است. این مسئله به دلیل پیچیدگی مؤثر بالاتر (وجود پارامترهای مؤثر بیشتر) مدل MVN است که با برازش بهتر جبران نمی‌گردد. در مقابل، گرچه مدل SC از منظر تعداد پارامترهای اثرات تصادفی پیچیده‌تر است اما به نظر می‌رسد اطلاعات را از راه مناسب‌تری اقتباس کند و بنابراین دارای تعداد مؤثر پارامترهای کمتر است. مدل SC توانایی بیشتری در تشخیص نواحی دارای خطر بالاتر دارد. در حالت وجود دو بیماری، مدل SC اندکی بهتر از مدل MVN عمل می‌کند. این مسئله ممکن است به دلیل فرض بسیار محدودکننده ثابت بودن همبستگی بین دو بیماری در تمام نواحی در مدل MVN باشد درحالی که مدل SC امکان متفاوت بودن خطر در نواحی مطالعه را از طریق پارامتر مقیاس فراهم می‌آورد (مک‌نب، ۲۰۰۹).

۵.۳ مقایسه مدل‌های SC و PM

برای هر دو مدل می‌توان کووریت‌های در سطح ناحیه ای را به پیشگویی خطی اضافه نمود. مدل PM از این لحاظ با مدل SC مشترک است که روندهای مشابه و غیرمشابه را کاملاً آشکار می‌کند (دابنی و همکاران، ۲۰۰۵).

۶.۳ مقایسه مدل‌های SC و PL

مدل SC برای نقشه بندی توأم بیماری‌ها طبیعی‌تر و منعطف‌تر از مدل PL است زیرا همه عامل خطر به عنوان اجزاء مشترک یا اختصاصی مدل در نظر گرفته می‌شوند. با این حال مدل PL دارای محاسن متعددی است. از آن جمله می‌توان به امکان تحلیل داده‌های مرگ و میر بدون دسترسی به اطلاعات جامعه ی در معرض خطر، امکان در نظر گرفتن پراکندگی در برآوردهای اثر سن در مدل، حذف میدان تصادفی نرمال مربوط به عبارات مشترک برای همه ی بیماری‌ها از مدل با استفاده از یک بیماری خاص به عنوان گروه مرجع و کاهش انحراف معیار برآوردهای عبارات ناهمگنی همبسته ی مکانی اختصاصی بیماری اشاره نمود. از سوی دیگر می‌توان از این نکته که انقباض اجزاء خوشه بندی در مدل PL نسبت به مدل SC بیشتر است به عنوان ایراد مدل PL یاد نمود (درسی، ۲۰۰۷).

۷.۳ مقایسه مدل‌های مکانی چندمتغیره و مکانی-زمانی چندمتغیره

تحلیل‌های مکانی-زمانی برتری‌هایی نسبت به تحلیل مکانی محض دارند زیرا امکان مطالعه ی هم‌زمان ثبات روندها در طول زمان و برجسته نمودن روندهای غیرمعمولی را فراهم می‌آورند. توانایی تعیین الگوهای مکانی خطرات بیماری‌ها که در طول زمان ثابت بوده یا افزایش می‌یابند، شواهد متقاعدکننده تری نسبت به تحلیل الگوهای مکانی در یک مقطع زمانی خاص ارائه می‌دهد و می‌تواند راهنمای مهمی برای تشریح علل بیماری و کمک به ایجاد خط مشی بهداشت محیطی فراهم آورد (اوسن و همکاران، ۲۰۰۸).

۴ بحث و نتیجه‌گیری

تاکنون مدل‌های گوناگونی برای نقشه‌بندی چندمتغیره‌ی بیماری‌ها مطرح شده است که هرکدام دارای مبانی تئوری، مزایا و معایب و نقاط قوت و ضعف خاص خود بوده و در مواردی خاص نسبت به دیگر مدل‌ها برتری دارند. در سال‌های اخیر مدل‌هایی معرفی شده‌اند که افزون بر تحلیل پراکندگی بیماری‌ها یا مرگ و میر ناشی از آن، الگوی پراکندگی مکانی عوامل خطر بیماری‌ها و نقش و اهمیت این عوامل برای بیماری‌های مرتبط را نیز آشکار کنند. استفاده از این مدل‌ها به ویژه مدل SC گسترش روزافزونی پیدا کرده و برتری آن نسبت به مدل‌هایی که تنها الگوی پراکندگی جغرافیایی بیماری‌ها را ارائه می‌دهند، نشان داده شده است. همچنین با توجه به اهمیت و کاربرد فراوان تحلیل روند زمانی میزان‌های بیماری‌ها، چند مدل نقشه‌بندی چندمتغیره‌ی مکانی-زمانی بیماری‌ها معرفی شده‌اند که نسبت به مدل‌های فاقد بعد زمانی دارای محاسن متعددی می‌باشند و تمرکز بر مطالعه و مقایسه‌ی این مدل‌ها نیز می‌تواند مفید و حائز اهمیت باشد. در این مقاله مجموعه‌ای از مهم‌ترین مدل‌های نقشه‌بندی چندمتغیره‌ی بیماری‌ها تشریح و مورد مقایسه قرار گرفته و تفاوت‌ها و شباهت‌های آن‌ها مورد بررسی قرار گرفت اما از آنجا که در سالیان اخیر تلاش‌های فراوانی در جهت گسترش و بسط بیش از پیش و روزافزون این مدل‌ها به عمل آمده است، لزوم مقایسه‌ی دامنه‌ی گسترده‌تری از مدل‌ها، هم از منظر تئوری و هم از منظر نتایج در یک مطالعه شبیه‌سازی جامع ضروری به نظر می‌رسد.

مراجع

- Ahmadipanah, M. V., Hassanzadeh, A. and Mahaki, B. (2018), Bivariate Spatiotemporal Shared Component Modeling: Mapping of Relative Death Risk Due to Colorectal and Stomach Cancers in Iran Provinces.
- Assunção, R. M. and Castro, M. S. (2004), Multiple Cancer Sites Incidence Rates Estimation Using a Multivariate Bayesian Model, *International journal of epidemiology*, **33**, 508-16.
- Besag, J., York, J. and Mollie, A., (1991), Bayesian Image Restoration with Two Applications In Spatial Statistics (with discussion), *Annals of the Institute of Statistical Mathematics*, **43**, 1-59.
- Böhning, D., (1999), Computer-Assisted Analysis of Mixtures and Applications: Meta-Analysis, Disease Mapping and Others, CRC Press; Taylor & Francis Group LLC, Florida, USA.
- Dabney, A. R. and Wakefield, J. C. (2005), Issues in the Mapping of Two Diseases, *Statistical Methods In Medical Research*, **14**, 83-112.
- Dreassi, E. (2007), Polytomous Disease Mapping to Detect Uncommon Risk Factors for Related Diseases, *Biometrical Journal: Journal of Mathematical Methods in Biosciences*, **49**, 520-9.
- Downing, A., Forman, D., Gilthorpe, M. S., Edwards, K. L. and Manda, S. O. (2008), Joint Disease Mapping Using Six Cancers in the Yorkshire Region of England, *International Journal of Health Geographics*, **7**, 41.

- Gelfand, A. E. and Vounatsou, P. (2003), Proper Multivariate Conditional Autoregressive Models for Spatial Data Analysis, *Biostatistics*, **4**, 11-5.
- Haddad, K. A., Jafari, K. T. and Mahaki, B. (2015), Investigating the Incidence of Prostate Cancer in Iran 2005–2008 Using Bayesian Spatial Ecological Regression Models, *Asian Pacific Journal of Cancer Prevention*, **16**, 5917-21.
- Held, L., Natário, I., Fenton, S. E., Rue, H. and Becker, N. (2005), Towards Joint Disease Mapping, *Statistical Methods in Medical Research*, **14**, 61-82.
- Jin, X., Carlin, B. P. and Banerjee, S. (2005), Generalized Hierarchical Multivariate CAR Models For Areal Data, *Biometrics*, **61**, 950-61.
- Khoshkar, A. H., Koshki, T. J. and Mahaki, B., (2015), Comparison of Bayesian Spatial Ecological Regression Models For Investigating the Incidence of Breast Cancer in Iran, 2005, *C*, **16**, 5669-73.
- Kim, H., Sun, D. and Tsutakawa, R. K. A. (2001), Bivariate Bayes Method for Improving the Estimates of Mortality Rates with a Twofold Conditional Autoregressive Model, *Journal of the American Statistical Association*, **96**, 1506-21.
- Knorr-Held, L. and Best, N. G. (2001), A Shared Component Model for Detecting Joint and Selective Clustering of Two Diseases, *Journal of the Royal Statistical Society: Series A*, **164**, 73–85.
- Langford, I. H., Leyland, A. H., Rasbash, J. and Goldstein, H. (1999), Multilevel Modelling of the Geographical Distributions of Diseases, *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, **48**, 253-68.
- Lawson, A. B., Browne, W. J. and Rodeiro, C. L. (2003), Disease Mapping with WinBUGS and MLwiN, John Wiley & Sons.
- Lawson, A. B., Biggeri, A. B., Boehning, D., Lesaffre, E., Viel, J. F. Clark, A. Schlattmann, P. and Divino, F. (2000), Disease Mapping Models: An Empirical Evaluation, Disease Mapping Collaborative Group, *Statistics in Medicine*, **19**, 2217-2241.
- Leyland AH, Langford IH, Rasbash J, Goldstein H, (2000 Sep 15), Multivariate spatial models for event data. *Statistics in Medicine*, **19**, 2469-78.
- MacNab, Y. C. (2009), Bayesian Multivariate Disease Mapping and Ecological Regression with Errors in Covariates: Bayesian Estimation of DALYs and 'Preventable' DALYs, *Statistics in Medicine*, **28**, 1369-85.
- Mahaki, B., Mehrabi, Y., Kavousi, A., Akbari, M. E., Waldhoer, T., Schmid, V. J. and Yaseri, M. (2011), Multivariate Disease Mapping of Seven Prevalent Cancers in Iran using Shared Component Model *Asian Pacific Journal of Cancer Prevention*, **12**, 2353.

- Mahaki, B., Mehrabi, Y., Kavousi, A. and Schmid, V. J. (2018), Joint Spatio-Temporal Shared Component Model with an Application in Iran Cancer Data, *Asian Pacific Journal of Cancer Prevention*, **19**, 1553-1560.
- Manda, S. O. and Leyland, A. (2070), An Empirical Comparison of Maximum Likelihood and Bayesian Estimation Methods for Multivariate Disease Mapping: Theory and Methods, *South African Statistical Journal*, **41**, 1-21.
- Marshall, R. J. (1991), Mapping Disease and Mortality Rates Using Empirical Bayes Estimators, *Applied Statistics*, **1**, 283-94.
- Mollie, A. (1990), Representation Geographique des Taux de Mortalire: Modelisation Spatiele et Methodes Bayesiennes, PhD thesis, Universite Paris 6, France.
- Nasrazadani, M., Maracy, M. R., Dreassi, E. and Mahaki, B. (2018), Mapping of Stomach, Colorectal, and Bladder Cancers in Iran, 2004–2009: Applying Bayesian polytomous Logit Model.
- Raei, M., Johann, S. V. and Mahaki, B. (2018), Bivariate Spatiotemporal Disease Mapping of Cancer of the Breast and Cervix Uteri Among Iranian Women, *it Geospatial Health*, **13**, 164-71.
- Rastaghi, S., Jafari, K. T. and Mahaki, B.,(2015), Application of Bayesian Multilevel Space Time Models to Study Relative Risk of Esophageal Cancer in Iran 2005 2007 at a County Level, *Asian Pacific Journal of Cancer Prevention*, **16**, 5787 92.
- Richardson, S., Abellan, J. J. and Best, N. (2006), Bayesian Spatio-Temporal Analysis of Joint Patterns of Male and Female Lung Cancer Risks in Yorkshire (UK), *Statistical Methods in Medical research*, **15**, 385-407.
- Simonoff, J. S.,(2012), Smoothing Methods in Statistics, Springer Science & Business Media.
- Tzala, E. H., Marshall, C. and Best, N. (2003), Multivariate Analysis Of Spatial And Temporal Variation In Cancer Mortality In Greece: Isee-398. *Epidemiology*, **14**, S79-80.
- Oleson, J. J., Smith, B. J. and Kim, H. (2008), Joint Spatio-Temporal Modeling of Low Incidence Cancers Sharing Common Risk Factors, *Journal of Data Science*, **6**, 105-23.



مدل سازی داده های فضایی نایستا

علی محمدیان مصمم، پریا عیوضی^۱
گروه آمار، دانشگاه زنجان

چکیده: در آمار فضایی اغلب برای سادگی فرض می شود که داده های موجود ایستا می باشند. لذا روش های مدل سازی داده ها و همچنین پیش بینی مشاهدات، مثلاً کریگینگ، اغلب بر اساس این فرض بنا شده اند. ولی فرض ایستایی داده های فضایی در عمل به دلیل تغییر پذیری داده ها به ویژه در مقیاس بزرگ دامنه فضایی برقرار نیست و بنابراین روش های موجود و کریگینگ کلاسیک برای این داده ها دقیق نخواهد بود. در این مقاله یک روش جدید برای پیش بینی داده های فضایی در حالت نایستایی مورد مطالعه قرار می گیرد. در این روش یک فضای ترکیبی جدید تعریف شده که این امر موجب می شود از مدل های فضایی- زمانی ایستا نظیر مدل تفکیک پذیر ضربی- جمعی و مدل تفکیک ناپذیر نایتینگ استفاده گردد. **واژه های کلیدی:** آمار فضایی، کریگینگ، فرآیند نایستا، مدل ضربی- جمعی، مدل نایتینگ **کد موضوع بندی ریاضی (۲۰۱۰):** 62M30.

۱ مقدمه

در آمار فضایی اغلب برای سادگی فرض می شود که داده های موجود ایستا باشند لذا روش های مدل سازی داده ها و همچنین پیش بینی مشاهدات، مثلاً کریگینگ، بر اساس این فرض بنا شده اند. ولی فرض ایستایی داده های فضایی در عمل به دلیل تغییر پذیری داده ها به ویژه در مقیاس بزرگ دامنه فضایی برقرار نیست و بنابراین روش های موجود و کریگینگ کلاسیک (کرسی، ۱۹۹۳) برای این داده ها دقیق نخواهد بود. زمین آمار (چیلز و دلفینر، ۲۰۰۹) امکان ترکیب انواع مختلف داده را فراهم می کند و از طریق استفاده از متغیرهای توضیحی (بیچامپ و همکاران، ۲۰۱۷) مناسب، فقدان ایستایی را در نظر می گیرد. مدل سازی روند خارجی یک روش مناسب برای انجام این کار است.

در این مدل یک ارتباط قطعی خطی با ضرایب مجهول بین متغیر برای برآورد و متغیر کمکی معرفی می شود و از آنجا که ضرایب روند، شناخته شده نیستند، این رابطه می تواند به طور محلی با استفاده از یک همسایگی متحرک تنظیم شود اما اگر روند دارای مقادیر بسیار پایینی باشد، برخی تخمین های غیر دقیق ممکن است رخ دهد. هدف این مقاله بررسی راه های

^۱ نام و ایمیل ارائه دهنده مقاله: پریا عیوضی، eivazi.paria@znu.ac.ir

مختلف در نظر گرفتن متغیر توضیحی بدون ایجاد فرض یک رابطه خطی بین داده‌ها و متغیر توضیحی است. در اینجا از متغیر توضیحی $\nu \in V$ به عنوان یک بعد جدید فضای ترکیبی $D \times V$ و $D \subseteq R^2$ استفاده می‌شود. فضای ترکیبی که در قبل معرفی شد برای ایجاد مجموعه‌ای از برآوردهای زمین‌آماري از جمله کریگینگ عادی، کریگینگ با روند خارجی، هم‌کریگینگ با یک مدل متریک یا ضربی-جمعی (دیاگو و همکاران، ۲۰۰۱، ۲۰۰۳) برای برازش مدل تغییرنگار استفاده می‌شود. با استفاده از تکنیک اعتبارسنجی متقابل^۱ عملکرد این برآوردها ارزیابی و مزایا و معایب آن‌ها مشخص می‌شود.

۲ روش کار

داده‌های فضایی $Z(s)$ جمع‌آوری شده از شبکه پایش کیفیت هوا را در نظر بگیرید. در کیفیت هوا، حتی اگر فرآیند اساسی مشاهدات جمع‌آوری شده از شبکه پایش نایستا باشند، $Z(s)$ گاهی به عنوان یک تابع تصادفی ایستای مرتبه دوم تلقی می‌شود و کریگینگ عادی برای پیش‌بینی داده‌ها در یک شبکه منظم با فرض ایستایی استفاده می‌شود. برآوردها Z^{OK} که از ترکیب خطی $\sum \lambda_\alpha Z_\alpha$ و وزن λ_α به دست آمده است، محدودیت‌های ناریبی و بهینگی را حل می‌کند و به سیستم کریگینگ عادی معروف است. برای پرداختن به مساله نایستایی، یک روش رایج مبتنی بر معرفی یک روند خطی $\mu(s) = \sum a_l f_l(s)$ به عنوان بخش قطعی $Z(s)$ است که مشاهدات را به خوبی توضیح می‌دهد:

$$Z(s) = \mu(s) + R(s).$$

با فرض اینکه نایستایی فرآیند در بخش قطعی محاسبه می‌شود، مدل تغییرنگار (یا کواریانس) برای بخش تصادفی $R(s)$ مورد نیاز است. در این چارچوب وزن‌های λ_α از حل سیستم کریگینگ با روند خارجی حاصل می‌شوند. باقیمانده R_α برای برآورد تغییرنگار نمونه مورد استفاده قرار می‌گیرد که به عنوان تفاوت بین مشاهدات واقعی Z_α و پیش‌بینی ارائه شده توسط روند برآوردها بدون در نظر گرفتن مکان s_α محاسبه می‌شود. در این حالت برآورد انجام شده بر روی باقیمانده‌های R_α تمایل به انحراف کمتری دارند. مدل‌سازی IRF-k^۲ روش دقیق برازش غیرمستقیم تغییرنگار است. ایده اصلی اینست که متغیر توضیحی $\nu \in V$ معمولاً به عنوان یک روند در زمین‌آمار دیده می‌شود و می‌توان آن را به عنوان یک مختصات فضایی معرفی کرد. در فضای ترکیبی $R^2 \times V$ ، $Z(s, \nu)$ به عنوان یک تابع تصادفی حقیقی مقدار تعریف می‌شود که با فرض اینکه ایستای مرتبه دوم است ν را مولفه جدیدی از تابع کواریانس در نظر می‌گیرد. تابع کواریانس به صورت زیر ارائه می‌شود:

$$C_{SV}(\|h_S\|, |h_V|) = Cov(Z(s + \|h_S\|, \nu + |h_V|), Z(s, \nu)).$$

۳ تعیین مدل تغییرنگار

در فضای ترکیبی $D \times V$ ، از فاصله اقلیدسی کلاسیک $\|h_S\|^2 + (k \cdot |h_V|)^2$ برای محاسبه تغییرنگار استفاده می‌شود.

$$\gamma_{SV}(\|h_S\|, |h_V|) = \gamma_m(\|h_{S \times V}\|)$$

در اینجا k پارامتر ناهمسان‌گردی مقیاس‌بندی بین R^2 و V است. بنابراین پس از تطبیق $\|h_S\|$ و $|h_V|$ و $\|h_{S \times V}\|$ از طریق k ، تمام فواصل مساوی در مدل موسوم به متریک γ_m بررسی می‌شوند. در این صورت تغییرنگارهای حاشیه‌ای $\gamma_{SV}(\cdot, |h_V|)$ و $\gamma_{SV}(\|h_S\|, \cdot)$ ویژگی‌های اساسی یکسانی دارند و در صورتی که مدل محدود باشد، آستانه یکسانی دارند.

^۱Cross-validation

^۲Intrinsic Random Fields of order k

به طور خاص، اگر آستانه γ_m به وسیله یکی از تغییرنگارهای حاشیه در فواصل کوچک به دست آید امکان اینکه γ_{SV} مقادیر بزرگتر از آستانه داشته باشد وجود ندارد، حتی اگر تغییرنگار نایستای تصادفی نمونه مقادیر بزرگتر را نشان دهد. استفاده از مدل جمعی-متریک به عنوان ترکیبی از مدل های فضایی، کمکی و متریک که لزوماً ساختار یکسانی را نشان نمی دهند می تواند به عنوان یک راهکار در نظر گرفته شود هر چند تطبیق داده ها با چنین ترکیبی دشوار خواهد بود. اگر کواریانس نایستای تصادفی به صورت تفکیک پذیر فرض شود می تواند به عنوان کواریانس های صرفاً فضایی و کمکی نمایش داده شود. اما به جای در نظر گرفتن مدل هایی با ویژگی های محدود می توان از خانواده های تفکیک ناپذیر (نایتینگ، ۲۰۰۲) عمومی استفاده کرد که در میان آن ها مدل ضربی-جمعی یک روش تفکیک ناپذیر انعطاف پذیر ارائه می دهد و قادر به تعامل با ناهمسانگردی های مختلف در فضای $D \times V$ می باشد. در اینجا یک روش یکسان برای جایگزین کردن کواریانس صرفاً زمانی $C_{T(h_T)}$ توسط یک کواریانس صرفاً کمکی $C_{V(h_V)}$ استفاده می شود:

$$C_{SV}(\|h_S\|, |h_V|) = k_1 C_S(h_S) C_V(h_V) + k_2 C_S(\|h_S\|) + k_3 C_V(|h_V|).$$

همچنین می توان از مدل نایتینگ به عنوان یک مدل دیگر از خانواده های تفکیک ناپذیر عمومی استفاده کرد که به صورت

$$C(\|h_S\|, |h_T|) = (\phi_T(|h_T|) + 1)^{-\frac{\delta}{2}} C_S(\|h_S\| (\phi_T(|h_T|) + 1)^{-\frac{1}{2}})$$

ارائه می شود، که در آن ϕ_T یک تابع دلخواه، C_S تابع کواریانس فضایی و $\delta > 0$ است. مدل ضربی-جمعی به دو دلیل نسبت به مدل نایتینگ برتری دارد:

- ۱- تغییرنگار نمونه اغلب داده ها شباهت بیشتری به رفتار تعریف شده توسط مدل ضربی-جمعی دارد.
- ۲- به نظر می رسد کواریانس نایستای تصادفی نمونه اغلب به صورت تفکیک پذیر با پارامتر k نزدیک به صفر پس از برازش روش بر روی یک مجموعه داده بزرگ شامل شود.

۴ تحلیل داده ها

در اینجا به ارزیابی مجموعه برآوردهای کریگینگ در یک مطالعه موردی برای میانگین روزانه دی اکسید نیتروژن در ۱۹ ژانویه ۲۰۱۱ می پردازیم. مطالعه انجام شده بر روی یک حوزه متمرکز در فرانسه است که در آن ایستگاه های مرتبط با ترافیک و صنعت مورد استفاده قرار نمی گیرند. برآوردهای مورد استفاده در این مطالعه با نماد Z_{vxx}^{kxx} ارائه می شود که در آن kxx برای نوع کریگینگ است و مقادیر آن را در حالت OK^3 ، KED^4 ، OCK^5 و $CKED^6$ می گیرد، همچنین vxx برای مدل تغییرنگار استفاده می شود یعنی زمانی که هیچ مقداری داده نشود یک تغییرنگار فضایی با فاصله اقلیدسی $\|h_S\|$ است درحالی که در فضای ترکیبی، اندیس های m و psm به ترتیب برای مدل های متریک و ضربی-جمعی مورد استفاده قرار می گیرد. ساختار مورد استفاده برای برازش تغییرنگار، مدل های معروف تغییرنگار هستند که عموماً در تحلیل کیفیت هوا استفاده می شوند، که شامل مدل کروی و نمایی برای تغییرنگار ایستا و مدل خطی برای در نظر گرفتن مولفه ذاتی در صورت نیاز، استفاده می شود. پیش بینی ها بر اساس فضای کلاسیک و فضای ترکیبی انجام گرفته است. برای مقایسه آن ها از معیار $RMSE = \sqrt{\frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2}$ استفاده شده است. نتایج پیش بینی ها در جدول ۱ ارائه شده است. با توجه به جدول ۱ روش Z_{psm}^{KED} دارای RMSE کمتری می باشد.

³ordinary kriging

⁴kriging with external drift

⁵ordinary cokriging

⁶cokriging with external drift

جدول ۱: جذر میانگین خطای پیش‌بینی برای روش‌های مختلف در حالت کلاسیک و فضای ترکیبی

نوع فضا	روش پیش‌بینی	RMSE
بدون استفاده از CHIMERE	Z^{OK}	۱۲/۸۴
	$Z^{KED}, f = (\nu)$	۱۱/۰۳
	Z_{psm}^{OK}	۱۰/۶۵
	Z_m^{OK}	۱۱/۱۵
	$Z^{KED}, f = (\varphi)$	۱۰/۹۷
با استفاده از CHIMERE	$Z^{KED}, f = (\varphi, \nu)$	۱۰/۵۶
	$Z_{psm}^{KED}, f = (\varphi)$	۹/۷۸
	$Z_m^{KED}, f = (\varphi)$	۱۱/۵۵
	Z^{OCK}	۱۳/۴۷
فضای ترکیبی	$Z^{CKED}, f = (\nu)$	۱۱/۱۸
	Z_{psm}^{OCK}	۹/۹
	Z_m^{OCK}	۱۰/۶۵

بحث و نتیجه‌گیری

تعریف نایستایی با در نظر گرفتن متغیر توضیحی به‌عنوان یک بعد از فضای ترکیبی می‌تواند نتایج کلاسیک با روند خارجی را بهبود بخشد. مطالعات روی داده‌های NO_2 نشان می‌دهد بهترین نتیجه با برازش مدل تغییرنگار ضربی-جمع‌ی روی باقیمانده‌های رگرسیون بین مشاهدات و متغیرهای توصیفی CHIMERE به‌دست آمده‌اند. در این مقاله برای سادگی مدل ضربی-جمع‌ی با رویکرد متریک مورد استفاده قرار گرفته است هرچند ممکن است که همیشه تضمین‌کننده بهترین مدل نباشد.

مراجع

- Beauchamp, M., de Fouquet, C., and Malherbe, L. (2017), Dealing with Non-Stationarity Through Explanatory Variables in Kriging-based Air Quality Maps, *Spatial Statistics*, **22**, 18-46.
- Chiles, J. P., and Delfiner, P. (2009), *Geostatistics: Modeling Spatial Uncertainty*, Vol. 497, John Wiley & Sons.
- Cressie, N. (1993), *Statistics for Spatial Data*, John Wiley, New York.
- De Iaco, S., Myers, D. E., and Posa, D. (2001), Space-time Analysis Using a General Product-Sum Model, *Statistics & Probability Letters*, **52**, 21-28.
- De Iaco, S., Myers, D. E., and Posa, D. (2003), The Linear Coregionalization Model and the Product-Sum Space-Time Variogram, *Mathematical Geology*, **35**, 25-38.
- Gneiting, T. (2002), Nonseparable, Stationary Covariance Functions for Space-Time Data, *Journal of the American Statistical Association*, **97**, 590-600.



تحلیل داده‌های فضایی-زمانی در حضور مقادیر گمشده

علی محمدیان مصمم، زهرا مقیمی^۱
گروه آمار، دانشگاه زنجان

چکیده: در بسیاری از تحقیقات با مواردی برخورد می‌کنیم که برای بعضی از متغیرها، پاسخ یا جوابی وجود ندارد یا اینکه پاسخ نامعلوم بوده و قادر به استخراج آن نیستیم. به این مقادیر، مقادیر نامعلوم یا داده‌های گمشده می‌گوییم. در داده‌های فضایی-زمانی با توجه به وجود وابستگی بین مشاهدات برحسب موقعیت فضایی و زمان مشاهده هر داده، مقادیر گمشده‌ای که در فواصل فضایی یا زمانی نزدیک‌تر نسبت به مشاهدات قرار دارند می‌توانند شامل اطلاعات مفیدی باشند که استفاده از آن‌ها می‌تواند منجر به نتایج دقیق‌تری شود. بنابراین لازم است وجود گمشدگی در داده‌ها مورد توجه و بررسی قرار گیرد. برای رسیدگی به مشکل داده‌های گمشده در این مقاله یک الگوریتم جدید مورد مطالعه قرار می‌گیرد. روش پیشنهادی هر مقدار گمشده را جداگانه براساس نقاط داده شده در یک محدوده فضایی-زمانی در اطراف آن نقطه پیش‌بینی می‌کند. پیش‌بینی مقادیر گمشده و تخمین عدم قطعیت پیش‌بینی مربوطه براساس روش‌های مرتب‌سازی و رگرسیون چندکی انجام گرفته است. در پایان روش پیشنهادی روی داده‌های حاصل از سنجش از راه دور مربوط به پوشش گیاهی منطقه‌ای از آلاسکا مورد استفاده قرار گرفته است.

واژه‌های کلیدی: داده‌های فضایی-زمانی، داده‌های گمشده، الگوریتم gap-fill، شاخص پوشش گیاهی.
کد موضوع بندی ریاضی (۲۰۱۰): 62M30.

۱ مقدمه

سنجش از دور^۱ یک تکنولوژی جدید برای مطالعه‌ی طیف گسترده‌ای از فرآیندهای سطح زمین از یک فاصله دور از زمین است. در این مقاله داده‌های پوشش گیاهی با استفاده از این تکنولوژی مورد تحلیل قرار می‌گیرد. این داده‌ها با توجه به اینکه در مقیاس مکانی و در طول زمان اندازه‌گیری می‌شوند، معمولاً دارای همبستگی فضایی-زمانی می‌باشند. یک مشکل اصلی در تحلیل چنین داده‌هایی وجود مقادیر زیادی داده‌های گمشده^۲ می‌باشد. به عنوان مثال، در اندازه‌گیری پوشش گیاهی سطح

¹ Remote Sensing

² Missing Values

^۱ نام و ایمیل ارائه دهنده مقاله: زهرا مقیمی، zahra_moghimi@znu.ac.ir

زمین توسط ماهواره‌ها به دلیل برخی موانع نظیر وجود ابر در مسیر ثبت پوشش گیاهی با حجم عظیمی از مقادیر گمشده مواجه می‌شویم. فقدان داده‌ها، معمولاً یکی از مسائل اصلی در حوزه‌ی جمع‌آوری داده می‌باشد، برای همین روش‌های مختلفی در برخورد با داده‌های گمشده ارائه گردیده است که برای جلوگیری از مشکلاتی مانند کمی حجم نمونه و یا عدم حضور واحدهایی از نمونه در مجموعه داده‌ها باید با استفاده از روش‌های مناسب جایگذاری شوند نظیر حذف داده گمشده از تحلیل، جایگذاری داده گمشده با مقادیر مناسب مانند میانگین، استفاده از برآوردهای رگرسیونی، روش داده افزایی، روش جانهی چندگانه که روشی مبتنی بر میانگین‌های شبیه‌سازی است (راگونتان، ۲۰۱۵). رابین (۱۹۷۶) انواع مختلف داده‌های گمشده را به سه دسته کاملاً تصادفی، تصادفی و غیر تصادفی طبقه‌بندی کرده است. دو استراتژی در این روش‌ها مشاهده می‌شود. استراتژی اول در برخورد با داده‌های گمشده استفاده از روش‌های آماری است که نسبت به حضور مقادیر گمشده مقاوم هستند. استراتژی دوم این است که مقادیر گمشده قبل از هرگونه تحلیل آماری مورد پیش‌بینی قرار گیرند هرچند کمتر مقاله‌ای برای داده‌های گمشده فضایی-زمانی می‌توان یافت. براساس استراتژی دوم در این مقاله یک الگوریتم جدید برای تخمین و جایگزینی داده‌های گمشده فضایی-زمانی مورد مطالعه قرار می‌گیرد. ادامه ساختار مقاله به این صورت می‌باشد که در بخش ۲ به معرفی الگوریتم gap-fill می‌پردازیم. در بخش ۳ به تحلیل داده‌های پوشش گیاهی با استفاده از الگوریتم مورد مطالعه پرداخته می‌شود. سرانجام بحث و نتیجه‌گیری در بخش ۴ ارائه شده است.

۲ معرفی الگوریتم gap-fill

یکی از روش‌های پیش‌بینی مقادیر گمشده فضایی-زمانی در مجموعه داده‌های سنجنش از راه دور استفاده از الگوریتم gap-fill است. این الگوریتم مشابه روش ارائه شده در وایس و همکاران (۲۰۱۴) است. با این تفاوت که روش جدید دارای کاربرد ساده‌تر بوده و علاوه بر این روش ما در انتخاب زیرمجموعه‌های همسایگی مقادیر گمشده باعث افزایش دقت می‌شود. الگوریتم gap-fill برای داده‌های ماهواره‌ای که عمدتاً دارای آرایه‌های چهار بعدی می‌باشند مناسب است. این چهار بعد به عنوان مثال می‌توانند شامل دو بعد فضایی برای طول و عرض جغرافیایی و دو بعد زمانی برای فصل و سال اندازه‌گیری‌ها باشد. از طرفی معمولاً داده‌های ماهواره‌ای داده‌های با اندازه بزرگ بوده لذا این الگوریتم باید دارای سرعت محاسباتی بالایی نیز باشد.

الگوریتم gap-fill (گربر و همکاران، ۲۰۱۸):

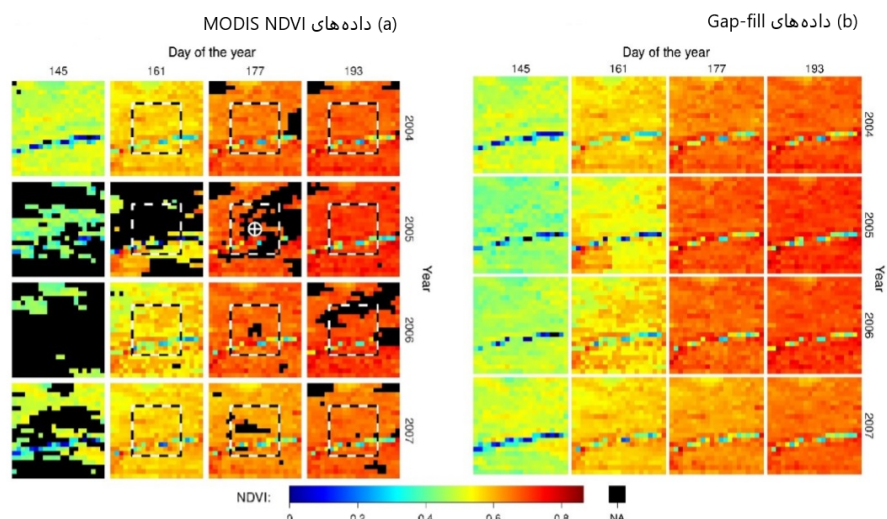
گام ۱: استخراج زیر مجموعه‌ای از داده‌ها

گام ۲: رتبه‌بندی تصاویر ماهواره‌ای

گام ۳: پیش‌بینی مقادیر گمشده براساس زیرمجموعه انتخابی با استفاده از رگرسیون چندکی^۳.

در ادامه گام‌های الگوریتم gap-fill برای یک داده‌ی گمشده تشریح می‌گردد و برای تمام مقادیر گمشده قابل تکرار است. در مرحله اول یک زیرمجموعه‌ای از داده‌های مورد نیاز پیش‌بینی از همسایگی مقدار گمشده انتخاب می‌شود. این زیرمجموعه خود شامل چهار بعد بوده که شامل دو بعد همسایگی فضایی و دو بعد همسایگی زمانی می‌باشد. توجه کنید که اندازه همسایگی‌ها طوری انتخاب می‌شود که فرآیند کلی برای مجموعه مشاهدات بهینه باشد. این کار معمولاً با استفاده از روش اعتبارسنجی متقابل انجام می‌گیرد. معیار انتخاب یک زیرمجموعه خوب این است که اولاً حداقل ۵ تصویر غیرگمشده در آن موجود باشد. ثانیاً اگر مقدار گمشده‌ای در زیرمجموعه انتخابی است حداقل ۲۵ تصویر گمشده در آن موجود باشد. در

³Quantile Regression



شکل ۱: داده‌های اصلی و داده‌های حاصل از الگوریتم gap-fill

مرحله دوم، تصاویر ماهواره‌ای براساس شاخص‌های پوشش گیاهی رتبه‌بندی می‌شوند. در مرحله سوم مقادیر گمشده براساس زیرمجموعه انتخابی با استفاده از رگرسیون چندکی خطی (کونیکر، ۲۰۰۵) پیش‌بینی می‌شوند. الگوریتم فوق در بسته آماری gap-fill در نرم‌افزار R موجود می‌باشد. از ویژگی مهم دیگر این الگوریتم این است که با محاسبات موازی می‌توان سرعت پردازش داده‌ها را افزایش داد. ما الگوریتم فوق را روی داده‌های MODIS NDVI بکار بردیم. توجه داشته باشید که روش پیشنهادی را می‌توان برای طیف وسیعی از داده‌های سنجش از راه دور اعمال کرد.

۳ تحلیل داده‌های پوشش گیاهی

شاخص نرمال شده تفاوت پوشش گیاهی^۴ که به اختصار NDVI نامیده می‌شود، شاخص گرافیکی ساده‌ای است که در تحلیل‌ها و اندازه‌گیری‌های سنجش از دور و ارزیابی وجود یا عدم وجود پوشش گیاهی یک منطقه کاربرد دارد. دامنه تغییرات این شاخص بین ۱+ و ۱- است. برای بررسی پوشش گیاهی هر منطقه‌ای می‌توان از روش‌های مختلفی استفاده کرد. اما بکارگیری تصاویر ماهواره‌ای می‌تواند تغییرات زمانی و مکانی پوشش گیاهی هر منطقه‌ای را با دقت بالاتری فراهم کند. در این زمینه پایگاه داده پوشش گیاهی سنجنده MODIS یکی از پایگاه داده‌های مهم است. یک مشخصه از داده‌های MODIS این است که مقدار ساختار وابستگی‌های فضایی-زمانی را از خود نشان می‌دهد. داده‌های مورد مطالعه در این مقاله دارای اندازه فضایی 21×21 پیکسل و شامل ۱۶ تصویر با ۴ شاخص فصلی (روزهای ۱۹۳، ۱۷۷، ۱۶۱، ۱۴۵ از هر سال) به صورت شبکه‌های منظم، مشاهده شده در طی ۴ سال (۲۰۰۴-۲۰۰۷) که آرایه داده در مجموع $7056 = 21 \times 21 \times 16$ عنصر دارد. خروجی الگوریتم بر روی داده‌ها در شکل ۱ ارائه شده است. شکل ۱(a) داده‌های اصلی پوشش گیاهی را نشان می‌دهد. این شکل شامل 21×21 پیکسل فضایی می‌باشد که در ۱۶ نقطه در ۴ روز متفاوت در یک سال است. داده‌های رسم شده شامل ۱۶۰۳ مقدار گمشده می‌باشد که بصورت بلوک‌های سیاه رنگ نمایش داده شده‌اند.

برای پیش‌بینی ۱۷۷ امین روز سال ۲۰۰۵ که با علامت سفید $\oplus$ رنگ نمایش داده شده است زیر مجموعه انتخابی با مستطیل خط چین مورد استفاده قرار گرفته است. در شکل ۱(b) خروجی gap-fill رسم شده است. نمودار فوق نشان

^۴Normalized Difference Vegetation Index

دهنده‌ی کارایی بسیار خوب الگوریتم مورد مطالعه است. مطالعات بعدی با استفاده از معیارهای مناسب آماری برای ارزیابی بهتر کارایی روش ارائه شده می‌باشد.

۴ بحث و نتیجه‌گیری

الگوریتم gap-fill یک الگوریتم انعطاف‌پذیر برای پیش‌بینی مقادیر گمشده داده‌های شاخص پوشش گیاهی می‌باشد. این الگوریتم به دلیل استفاده از محاسبات موازی و روش‌های پیش‌بینی آماری دارای سرعت و دقت بهتری نسبت به الگوریتم‌های موجود در این زمینه می‌باشد. از مزایای دیگر این الگوریتم این است که نیازی به فرض نرمال بودن داده‌ها ندارد. علاوه بر این، این الگوریتم دارای کارایی بالایی در حضور مقادیر گمشده با حجم زیاد می‌باشد. مهمترین مشکل آن می‌تواند برای حالتی باشد که داده‌ها روی شبکه منظم قرار نگرفته باشند.

مراجع

- Gerber, F., de Jong, R., Schaepman, M. E., Schaepman-Strub, G., & Furrer, R. (2018), Predicting Missing Values in Spatio-Temporal Remote Sensing Data, *IEEE Transactions on Geoscience and Remote Sensing*, **56**(5), 2841-2853.
- Koenker, R. (2005), *Quantile Regression (Econometric Society Monographs)*, Cambridge: Cambridge University Press.
- Raghunathan, T. (2015), *Missing Data Analysis in Practice*, Chapman and Hall/CRC.
- Rubin, D. B. (1976), Inference and Missing Data. *Biometrika*, **63**(3), 581-592.
- Weiss, D. J., Atkinson, P. M., Bhatt, S., Mappin, B., Hay, S. I., & Gething, P. W. (2014), An Effective Approach for Gap-filling Continental Scale Remotely Sensed Time-series, *ISPRS Journal of Photogrammetry and Remote Sensing*, **98**, 106-118.



تحلیل بیزی داده‌های گسسته فضایی-زمانی

علی محمدیان مصمم، الناز عباسی^۱
گروه آمار، دانشگاه زنجان

چکیده: در تحلیل بیزی داده‌ها گاهی با داده‌های گسسته‌ای مواجه می‌شویم که به دلیل ناگوسی بودن توزیع متغیر پاسخ و وجود تعداد زیادی متغیر پنهان در مدل تحت بررسی شکل بسته‌ای برای توزیع پسینی وجود ندارد. از طرفی ممکن است متغیرهای پاسخ در طول زمان به یکدیگر وابسته باشند. در این شرایط به دلیل عدم استقلال داده‌ها یا ناگوسی بودن توزیع داده‌ها استفاده از روش‌های مونت کارلوی زنجیر مارکف ($MCMC$) با چالش‌هایی نظیر، وجود پارامترهای زیاد در ساختار سلسله مراتبی، محاسبات سنگین، شبیه‌سازی گسترده، طولانی و زمان‌بر بودن محاسبات به‌ویژه زمانی که بعد میدان تصادفی بزرگ است و عدم همگرایی توزیع پسینی مواجه هستیم. برای حل این مشکلات (رو و همکاران، ۲۰۰۹) روش تقریب لاپلاس آشیانی جمع بسته ($INLA$) را پیشنهاد دادند. این روش ضمن حفظ دقت برآوردها و پیش‌گویی‌های روش $MCMC$ سرعت محاسبات قابل قبولی دارد. در این مقاله در یک مطالعه موردی به تحلیل داده‌های جرم و جنایت کشور کانادا با رهیافت $INLA$ می‌پردازیم.

واژه‌های کلیدی: تقریب لاپلاس آشیانی جمع بسته، آمار بیزی، آمار فضایی-زمانی

کد موضوع بندی ریاضی (۲۰۱۰): 62M30.

۱ مقدمه

تحلیل بیزی داده‌های جرم و جنایت معمولاً به صورت استنباط بیزی الگوهای فضایی محض یا الگوهای زمانی محض انجام می‌گیرد. اگر مشاهدات به‌گونه‌ای باشند که مشاهدات نزدیک به هم وابسته‌تر و مشاهدات دورتر وابستگی کمتری داشته باشند، این‌گونه مشاهدات داده‌های فضایی نامیده می‌شوند. بدیهی است که تحلیل آماری داده‌های فضایی با روش‌های آماری معمول مقدور نیست، زیرا شرط اساسی استقلال داده‌ها تحقق پیدا نمی‌کند. با این حال درمورد داده‌های جرم و جنایت چنین تحلیل‌های بیزی صرفاً فضایی یا زمانی به دلیل اینکه اثر متقابل فضا و زمان را در نظر نمی‌گیرند مدل‌های مناسبی نیستند. در این مقاله سعی شده است که با استفاده از مدل‌های بیزی فضایی-زمانی به تحلیل و تفسیر این اثر نیز پرداخته شود.

^۱ نام و ایمیل ارائه دهنده مقاله: الناز عباسی، Elnaz.Abbasi@Znu.ac.ir

مشاهداتی که هم از نظر موقعیت فضایی و هم از نظر موقعیت زمانی وابسته باشند، داده‌های فضایی-زمانی نامیده می‌شوند. در آمار فضایی-زمانی به طور معمول یک میدان تصادفی به عنوان مدل آماری داده‌های فضایی-زمانی در نظر گرفته می‌شود. میدان تصادفی فضایی-زمانی مجموعه‌ای از متغیرهای تصادفی مانند $Z(s, t); (s, t) \in D \times T$ است، که در آن s موقعیت فضایی در مجموعه $D \subseteq R^d$ ، $d \geq 1$ و t لحظه زمانی در مجموعه $T \subseteq R$ است.

روش مرسوم برای محاسبات بیزی استفاده از روش‌های مونت کارلوی زنجیر مارکف (MCMC) مشکلات فراوانی به دلایلی نظیر وابستگی شدید بین مشاهدات و بعد زیاد ابرپارامترها در محاسبات توزیع‌های پسینی را تجربه کرده است. لذا هدف این مقاله استفاده از روش تقریب لاپلاس آشیانی جمع بسته (INLA) برای تحلیل چنین داده‌هایی است. روش INLA اولین بار توسط (رو و همکاران، ۲۰۰۹) ارائه شد و سپس توسط (مارتین و همکاران، ۲۰۱۳) و (رو و همکاران، ۲۰۱۷) توسعه پیدا کرد.

ادامه ساختار مقاله به صورت زیر می‌باشد. در بخش ۲ روش INLA معرفی می‌شود. در بخش ۳ داده‌های تماس به پلیس برای گزارش جرم و جنایت مدل بندی می‌شوند، پارامترهای مدل برآورده شده و با استفاده از معیار بیزی اطلاع انحراف DIC بهترین مدل انتخاب می‌شود.

۲ تقریب لاپلاس آشیانی جمع بسته

همانگونه که در مقدمه ملاحظه گردید روش‌های عددی مرسوم در آمار بیزی نظیر روش MCMC دارای مشکلات اساسی به ویژه سرعت همگرایی و مدت زمان لازم برای محاسبات عددی می‌باشند. در این بخش به معرفی روش جایگزین، INLA، می‌پردازیم. INLA یک رهیافت جدید برای استنباط تقریبی در مدل‌های گاوسی پنهان (LGM) است. مطالعات نشان می‌دهند که برای چنین مدل‌هایی اغلب اوقات روش INLA سریع‌تر و دقیق‌تر از روش MCMC می‌باشد. چهارچوب روش INLA نیازمند سه مؤلفه کلیدی است: مدل‌های گاوسی پنهان، میدان تصادفی مارکف گاوسی (GMRF) (لینگرین و همکاران، ۲۰۱۱) و تقریب لاپلاس می‌باشد. اغلب مدل‌های بیزی در رگرسیون، آمار فضایی و آمار فضایی-زمانی دارای ساختار گاوسی پنهان می‌باشند. علاوه بر این در کاربست INLA نیازمند یک میدان تصادفی مارکف گاوسی می‌باشیم. GMRF یک میدان تصادفی گاوسی با خاصیت استقلال شرطی است یعنی x و x' به شرط سایر مؤلفه‌ها مستقل هستند اگر و تنها اگر i و j زامین مؤلفه ماتریس دقت Q صفر هستند.

برای توصیف مفاهیم فوق مدل سلسله مراتبی سه مرحله‌ای زیر را در نظر بگیرید، فرض کنید که در مرحله اول متغیر پاسخ $y = (y_1, \dots, y_n)'$ به‌طور شرطی مستقل و دارای توزیع نمایی خاصی باشد

$$y|x, \theta_1 \sim \prod_{i=1}^n P(y_i|x_i; \theta_1). \quad (1.2)$$

در مرحله دوم x را یک میدان تصادفی گاوسی پنهان با تابع چگالی

$$P(x|\theta_2) \propto |Q_{\theta_2}|_+^{1/2} \exp(-1/2 x' Q_{\theta_2} x) \quad (2.2)$$

در نظر بگیرید که در آن ماتریس دقت Q_{θ_2} یک ماتریس معین مثبت با پارامتر θ_2 و $|Q_{\theta_2}|_+$ حاصل ضرب مقادیر ویژه غیر صفر آن می‌باشد و وارون همان ماتریس واریانس کواریانس می‌باشد. در مرحله نهایی فرض کنید که پارامتر $\theta = (\theta_1, \theta_2)$ دارای یک توزیع پیشینی $\pi(\theta)$ است. در نتیجه پسینی پارامترهای مدل سلسله مراتبی به صورت

$$\begin{aligned} \pi(\eta, \theta|y) &\propto \pi(\theta) \pi(\eta|\theta_2) \prod_{i=1}^n \pi(y_i|\eta_i; \theta_1) \\ &\propto \pi(\theta) |Q_{\theta_2}|_+^{1/2} \exp\{-1/2 \eta' Q_{\theta_2} \eta + \sum_{i=1}^n \log P(y_i|\eta_i; \theta_1)\} \end{aligned} \quad (3.2)$$

می‌باشد. توزیع‌های پسینی حاشیه‌ای برای متغیرهای پنهان و ابر پارامترها به صورت

$$\begin{aligned}\pi(\eta_i|y) &= \int \pi(\eta_i|\theta, y)\pi(\theta|y)d\theta \\ \pi(\theta_i|y) &= \int \pi(\theta|y)d\theta_{-i}\end{aligned}\quad (۴.۲)$$

به دست می‌آیند. (رو و همکاران، ۲۰۱۷) یک تقریب لاپلاس برای توزیع‌های حاشیه‌ای پسینی ارائه کردند. با استفاده از روش تقریب لاپلاس در *INLA* توزیع‌های پسینی هر یک از پارامترهای نامعلوم به دست می‌آید.

۳ مدل بندی تماس برای خدمات پلیس

منطقه آنتاریو در ایالت واترلو شامل شهرهای واترلو، کیچینر، کمبریج و ۴ منطقه روستایی می‌باشد. جمعیت این منطقه در سال ۲۰۱۱، ۵۰۶۱۰۷ نفر بوده که در ۷۵۵ ناحیه‌ی سرشماری توزیع شده است. در کل ۲۹۰۲۷ تماس به پلیس برای گزارش جرم و جنایت در طول ۹۰۶۰ واحد فضا- زمان ثبت شده است (لوان و همکاران، ۲۰۱۶). این داده‌ها برای سال ۲۰۱۱ در ۱۲ دوره ۲ ساعته طبقه‌بندی شده‌اند. در این مقاله برای مدل‌سازی فضایی-زمانی بیزی از مدل سلسله مراتبی سه مرحله‌ای و رهیافت *INLA* استفاده می‌شود. در مرحله اول چون تعداد تماس‌های جرم و جنایت کم می‌باشد لذا تعداد تماس‌ها برای خدمات پلیس در منطقه i ، ۷۵۵، $i = 1, \dots, 755$ و دوره زمانی t ، ۱۲، $t = 1, \dots, 12$ دارای توزیع پواسن به صورت

$$y_{it}|\mu_{it} \sim POI(\mu_{it}) \quad i = 1, \dots, 755, \quad t = 1, \dots, 12 \quad (۱.۳)$$

فرض شده است. میانگین توزیع تعداد تماس‌ها را به صورت

$$\log(\mu_{it}) = \log(E_{it}) + \alpha + u_i + s_i + \gamma_t + \phi_t + \psi_{it} \quad (۲.۳)$$

در نظر بگیرید که در آن E_{it} تعداد مورد انتظار تماس‌ها بوده و ریسک نسبی به ریسک کلی یا اثر ثابت (α) ، اثرات تصادفی فضایی $(u_i + s_i)$ ، اثرات تصادفی زمانی $(\gamma_t + \phi_t)$ و اثر تصادفی فضایی-زمانی (ψ_{it}) تجزیه شده است. در مرحله دوم فرض کنید

$$\begin{aligned}u_i &\sim Normal(0, \sigma_u^2), \quad i = 1, \dots, 755 \\ s_i &\sim ICAR(W, \sigma_s^2) \\ \gamma_t &\sim Normal(0, \sigma_\gamma^2), \quad t = 1, \dots, 12 \\ \phi_t &\sim ICAR(P, \sigma_\phi^2)\end{aligned}\quad (۳.۳)$$

که در آن σ_ϕ^2 و σ_γ^2 ابر پارامترها می‌باشند. همان‌گونه که ملاحظه می‌گردد اثرات تصادفی فضایی و اثرات تصادفی زمانی به دو مؤلفه بدون ساختار و ساختاریافته تجزیه شده‌اند که مدل‌های ساختاریافته شامل مدل *ICAR* می‌باشند که به ترتیب همبستگی فضایی و زمانی این اثرات را مدل بندی می‌کنند. *ICAR* یک توزیع پیشینی مرسوم در آمار فضایی برای پارامترهای اثرات تصادفی می‌باشد که در آن میانگین مورد انتظار s_i برابر میانگین همسایگی آن می‌باشد. برای مرحله سوم تحلیل بیزی سلسله مراتبی توزیع پیشینی تمام ابر پارامترها $Gamma(0.5, 0.0005)$ در نظر می‌گیریم. برای تحلیل بیزی سلسله مراتبی از روش *INLA* استفاده می‌کنیم. چهار مدل متفاوت برای تحلیل بیزی داده‌ها در نظر گرفته شده است: (۱) ψ_{it} ‌ها مستقل و هم توزیع هستند.

(۲) ψ_{it} ها دارای همبستگی فضایی و ناهمبستگی زمانی هستند.

(۳) ψ_{it} ها دارای ناهمبستگی فضایی و همبستگی زمانی هستند.

(۴) ψ_{it} ها دارای همبستگی فضایی و زمانی هستند.

۱.۳ انتخاب مدل

مدل سازی در نرم افزار R با روش $INLA$ با ۴ مدل فوق انجام شده است. ورودی های تابع $INLA$ و نقشه ها با نرم افزار GIS تهیه شده اند. DIC برای انتخاب مدل بهتر از معیار DIC استفاده می شود. DIC معیار بیزی برای نیکویی برازش و تعیین پیچیدگی مدل است. (اشپگل و همکاران، ۲۰۰۳) پیشنهاد دادند که امید ریاضی پسین آماره انحراف $\bar{D}$ و تعداد پارامترهای مؤثر در مدل p_D اطلاعات مدل را خلاصه می کنند. بنابراین معیار اطلاع انحراف DIC به صورت $DIC = \bar{D} + p_D$ است: مدلی که کمترین DIC را نسبت به بقیه مدل ها داشته باشد مدل بهتری است. مقادیر DIC و تعداد پارامترهای مؤثر برای ۴ مدل فضایی-زمانی و تعامل فضا-زمان در جدول ۱ آورده شده است. طبق جدول ۱.۳ مدل اول با کمترین DIC بهترین مدل انتخاب شده است.

جدول ۱: مقایسه DIC مدل ها

تعامل فضا-زمان	پارامترهای تعامل	اطلاع انحراف	پارامتر مؤثر
مدل اول	γ_t و u_i	۵۴۰۸۹	۴۹۸۰
مدل دوم	ϕ_t و u_i	۷۲۷۹۱/۲۸	۷۵۸
مدل سوم	γ_t و s_i	۵۴۴۹۶	۴۸۰۰/۷۲
مدل چهارم	ϕ_t و s_i	۷۲۷۹۱/۵۲	۷۵۸/۴۴

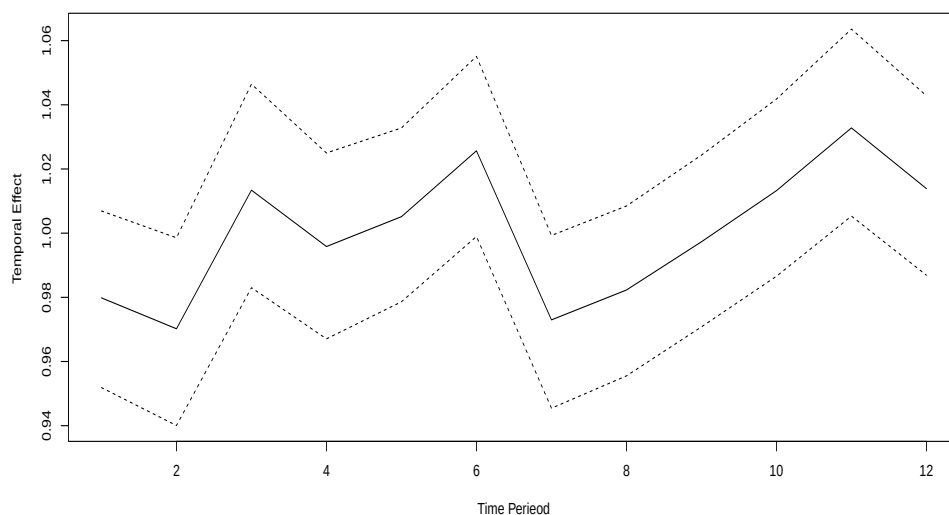
۲.۳ برآورد پارامترهای مدل

برآورد پارامترهای مدل اول با روش $INLA$ در جدول ۲ آورده شده است. نمودار ۱ نشان دهنده اثرات تصادفی زمانی

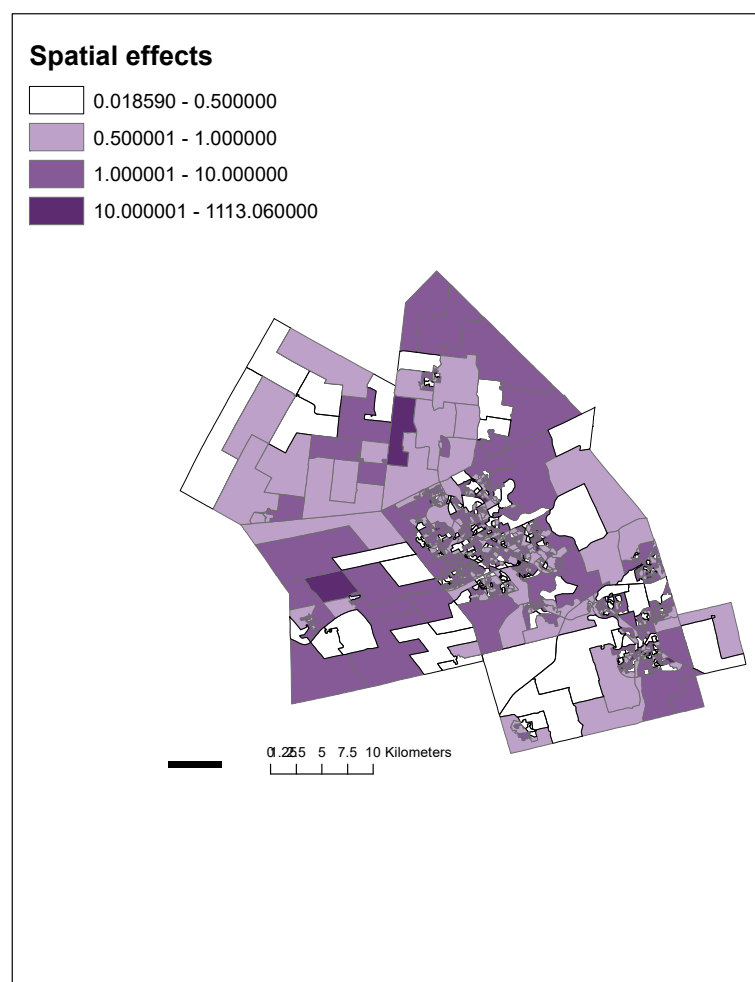
جدول ۲: برآورد پارامترهای مدل اول

پارامترهای مدل	میانگین	انحراف پسینی	چارک اول	میانه	چارک سوم	مد
σ_u^2	۰/۹۱۹	۰/۴۸۹	۰/۸۲۷	۰/۹۱۷	۱/۰۱۹	۰/۹۱۴
σ_γ^2	۱۷۸۹/۲۹۰	۹۸۷/۱۵۱	۵۴۴/۹۷۶	۱۵۷۴/۷۰۸	۴۲۸۹/۶۶۲	۱۲۰۲/۱۶۲
σ_ψ^2	۱۳/۰۳۴	۰/۳۵۵	۱۲/۳۱۵	۱۳/۰۲۸	۱۳/۷۱۳	۱۳/۰۵۴

است. براساس این نمودار بیشترین تأثیرات اصلی زمانی در دوره زمانی بین ۵۹ : ۱۱ - ۱۰ : ۰۰ و ۵۹ : ۲۱ - ۲۰ : ۰۰ است. بنابراین تخصیص منابع انسانی پلیس، باید در این دوره زمانی بیشتر باشد. شکل ۲ نشان دهنده اثرات تصادفی فضایی است. مناطقی با اثرات فضایی بالا، به طور مداوم تعداد زیادی تماس خدمات دارند و مناطقی هستند که منابع پلیس باید برای تمام دوره های زمانی به کار گرفته شود. در حالت کلی اثرات اصلی فضایی در مراکز منطقه مورد مطالعه بیشترین و در مناطق اطراف روستایی به ویژه در شمال غربی کمترین میزان است.



شکل ۱: اثرات تصادفی زمانی همراه با فاصله باورمندی ۹۵ درصد



شکل ۲: اثرات تصادفی فضایی

۴ بحث و نتیجه‌گیری

در این مقاله، در یک مطالعه موردی به تحلیل بیزی داده‌های جرم و جنایت کشور کانادا با رهیافت *INLA* پرداخته شد. برای این منظور چهار مدل مختلف به داده‌ها برازش و بهترین مدل با استفاده از معیار *DIC* انتخاب گردید.

مراجع

- Luan, H., Quick, M., and Law, J. (2016), Analyzing 1 Local Spatio-Temporal Patterns of Police Calls-for-Service Using Bayesian Integrated Nested Laplace Approximation, *International Journal of Geo-Information*, **5**(9), .162
- Lindgren, F., Rue, H., and Lindström, J. (2011), An Explicit Link Between Gaussian Fields and Gaussian Markov Random Fields: The Stochastic Partial Differential Equation Approach, *Journal of the Royal Statistical Society: Series B*, **73**(4), 423-498.
- Martins, T. G., Simpson, D., Lindgren, F., and Rue, H. (2013), *Bayesian computing with INLA: new features*. Computational Statistics and Data Analysis, **67**, .68-83
- Rue, H., Martino, S., and Chopin, N. (2009), Approximate Bayesian Inference for Latent Gaussian Models by Using Integrated Nested Laplace Approximations, *Journal of the Royal Statistical Society: Series B*, **71**(2), .319-392
- Rue, H., Riebler, A., Sørbye, S. H., Illian, J. B., Simpson, D. P., and Lindgren, F. K. (2017), Bayesian Computing with INLA: A Review, *Annual Review of Statistics and Its Application*, **4**, 395-421.
- Spiegelhalter, D., Best, N. G., Carlin, B. P., and van der Linde, A. (2003), Bayesian Measures of Model Complexity and Fit, *Quality Control and Applied Statistics*, **48**(4), .431-432



پایش کیفیت آب‌های سطحی با استفاده از تکنیک‌های آمار فضایی: مطالعه موردی استان زنجان

سرور موسوی^۱، یونس خسروی، عباسعلی زمان، مهسا دادشی
گروه علوم محیط زیست، دانشگاه زنجان

چکیده: آب‌های سطحی جاری یا رودخانه‌ها از مهم‌ترین منابع آب هستند که در تامین آب مورد نیاز فعالیت‌های مختلف مانند کشاورزی، صنعت، شرب و تولید برق نقش مهمی دارند. برخی کشورها منابع آب سطحی خود را براساس توسعه پایدار برنامه‌ریزی می‌کنند. یکی از نیازمندی‌های مهم در برنامه‌ریزی و توسعه منابع آب و حفاظت و کنترل آن‌ها آگاهی از کیفیت منابع آب است. کیفیت و آلودگی آب‌های سطحی تحت تاثیر تغییرات طبیعی و انسانی قرار دارند. در این مطالعه ارزیابی کیفیت آب‌های سطحی استان زنجان از نظر آلاینده ترکیب بیکربنات برای ۲۰ سال نمونه‌برداری شد. برای درک بهتر فرایندهای شیمیایی و بررسی کیفیت آب از روش‌های آماری فضایی هم‌چون شاخص‌های موران محلی و لکه‌های داغ و نرم‌افزار ArcGIS استفاده شده است. این مطالعه با هدف شناسایی عامل‌های موثر بر میزان و پراکنش بیکربنات انجام گرفته است. نتایج نشان داد که هیچ لکه‌ی داغی در روش لکه‌های داغ وجود ندارد. اما روش موران محلی حاکی از خوشه‌بندی بالا-بالا و پایین-پایین در مقادیر بیکربنات است. که احتمال می‌رود متاثر از نوع زمین‌شناسی و کاربری اراضی منطقه مورد مطالعه باشد. به منظور مدیریت بهتر محیط، می‌توان ایستگاه‌های مهم منبع آلودگی در آب‌های سطحی استان زنجان را شناسایی کرد. شاخص‌های موران محلی و لکه‌های داغ ابزاری مفید برای شناسایی نقاط مهم آلودگی در آب‌های سطحی استان زنجان است.

واژه‌های کلیدی: آب‌های سطحی، آمار فضایی، شاخص موران محلی، شاخص لکه‌های داغ
کد موضوع‌بندی ریاضی (۲۰۱۰): 62M30.

۱ مقدمه

مشکل کمبود آب و کیفیت آن از دیدگاه سازمان ملل متحد پس از مسئله جمعیت به عنوان دومین مشکل اصلی جهان شناخته شده است. با ورود به هزاره سوم و همزمان با افزایش جمعیت به ویژه در کشورهای در حال توسعه، تقاضا

^۱ نام و ایمیل ارائه دهنده مقاله: سرور موسوی، s.musavi@Znu.ac.ir

برای آب به منظور تأمین نیازهای جمعیتی افزایش قابل ملاحظه‌ای یافته است و سبب شده تا استفاده بهینه از آب و حوضه‌های آبی موجود در جهان با تأمل و برنامه‌ریزی بهتری حرکت کند (محمدجانی و یزدانیان، ۲۰۱۴). امروزه افزایش آلاینده‌های شیمیایی شهری و صنعتی و روش‌های نوین کشاورزی تهدیدی جدی برای محیط‌زیست به حساب می‌آیند. برای ارزیابی توسعه‌ی یک منطقه، کیفیت آب از مهم‌ترین مولفه‌ها است. کیفیت آب براساس متغیرهای فیزیکی و شیمیایی تعریف می‌شود. ترکیب شیمیایی آب سطحی، به عنوان منبع آبی برای مصرف انسانی و حیوانی، آبیاری و صنعتی را شامل می‌شود. بنابراین هدف، تعریف کیفیت آب نیست، بلکه استفاده مطلوب از آب در جامعه است (دشتی برمکی و همکاران، ۱۳۹۳) (فرچیچی و همکاران، ۲۰۱۸). یکی از نیازمندی‌های مهم در برنامه‌ریزی و توسعه منابع آب، حفاظت و کنترل آن‌ها و آگاهی از کیفیت منابع آب است. موردهایی هم‌چون انتشار مواد مغذی کشاورزی و مواد زاید شهری، افزایش تقاضای آب، استانداردهای سطح بالای زندگی و کاهش منابع قابل قبول آب سبب آلودگی آب و ایجاد وضعیت نامناسب محیط‌زیستی در سراسر جهان شده است. از این رو درک بهتر تغییرات آلودگی سیستم‌های آبی ضرورت بیش‌تری یافته است (فریادی و همکاران، ۱۳۹۱). رودخانه‌ها جز کوچکی از منابع آب‌های جهان می‌باشند که به عنوان یکی از منابع اساسی تأمین آب به حساب می‌آیند (لی و همکاران، ۲۰۱۸). امروزه تحلیل‌های آمار فضایی به عنوان ابزارهای جدید و کارآمد در علوم مختلف مورد توجه محققین قرار گرفته‌اند. توجه به عدم استقلال مشاهدات و وابسته بودن آن‌ها به یکدیگر در فضای مورد مطالعه و استفاده مستقیم از فضا، محیط، همسایگی، جهت‌گیری و روابط فضایی در محاسبات سبب برتری یافتن این روش نسبت به آمار کلاسیک شده است. در مطالعه‌ی حاضر روابط مکانی داده‌های آب‌های سطحی استان زنجان از نظر مقدار بیکربنات با استفاده از آمار فضایی بررسی شده است. به عنوان یکی از مهم‌ترین هدف‌های آمار فضایی برای تحلیل الگوهای فضایی و درک وابستگی‌های فضایی است که با استفاده از شاخص‌های موران محلی و لکه‌های داغ صورت می‌گیرد. (بحری و خسروی، ۱۳۹۶).

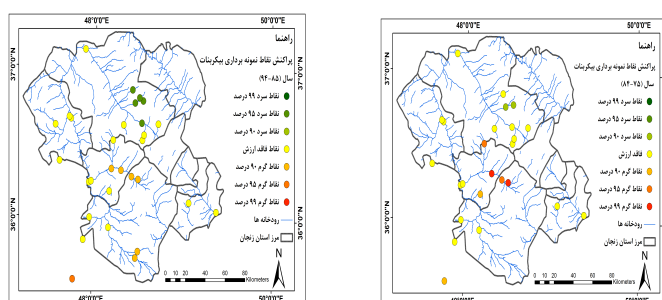
۲ مواد و روش‌ها

استان زنجان با وسعتی بیش از ۲۲ هزار کیلومتر مربع در منطقه شمال غربی کشور بین ۳۵ درجه و ۳۳ دقیقه الی ۳۷ درجه و ۱۵ دقیقه و ۳۷ درجه عرض شمالی ۴۷ درجه و ۲۶ دقیقه طول شرقی از نصف‌النهار قرار دارد. در این مطالعه از رودهای اصلی استان زنجان شامل زنجان‌رود، سجاس‌رود، رود شورچای، رودایجرود، رود قزل‌اوزن، رود ابهررود و رود خرارود و علاوه بر رودهای ذکر شده از رودهای فرعی دیگری که در استان زنجان جریان دارند نمونه‌برداری شده است. پارامتر مورد مطالعه بیکربنات است که یکی از ترکیب‌های شیمیایی تاثیرگذار بر سلامت آب شرب است. مقدار بیکربنات کمتر از ۹۰ ppm برای کشاورزی مناسب است و بیش از آن آبیاری را دچار مشکل می‌کند. بنابراین از ایستگاه‌های مختلف استان زنجان در طی ۲۰ سال نمونه‌برداری صورت گرفته است که به دو گروه ۱۰ ساله اول (از سال ۱۳۷۵ تا ۱۳۸۴) و ۱۰ ساله دوم (از سال ۱۳۸۵ تا ۱۳۹۴) تفکیک شده است. داده‌های طبقه‌بندی شده با استفاده از شاخص موران محلی (*Local Moran*) و شاخص لکه‌های داغ (*Hotspot*) و نرم‌افزار *ArcGIS* مورد بررسی قرار گرفته شد.

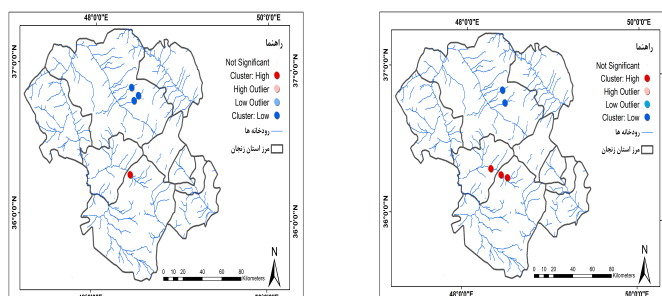
۳ بحث و نتیجه‌گیری

براساس نتایج بدست آمده، عدد خروجی شاخص لکه‌های داغ و عدد *Z* بدست آمده برای بیکربنات در ده سال اول (۰/۳۱) و برای ده سال دوم (۰/۷۳ -) تایید می‌کند که هیچ‌گونه خوشه‌بندی در مقدار بیکربنات در آب‌های سطحی وجود ندارد. خروجی شاخص موران محلی نشان داد با عدد *Z* برابر با ۲/۱۷ در سطح اطمینان ۹۵٪ برای ده سال اول (۷۵-۸۴) خوشه‌بندی در مقدارهای بالا-بالا را تشکیل داده است. هم‌چنین عدد *Z* برابر با ۲/۶۸ در سطح اطمینان ۹۹٪ برای ده

سال دوم (۸۵-۹۴) نشان داد که خوشه‌بندی در مقادیر بالا-بالا تشکیل شده است. از نظر شاخص لکه‌های داغ در ده سال اول در سطح اطمینان ۹۰٪ نقاط سردی در رودخانه زنجان رود شهر زنجان مشاهده می‌شود و نقاط گرم در سطح اطمینان ۹۵٪ و ۹۹٪ نیز در رودخانه سجاس رود شهر ایجرود مشاهده می‌شود. هم‌چنین در ده سال دوم نیز نقاط سرد در سطح اطمینان ۹۵٪ در رودخانه زنجان رود شهر زنجان و نقاط گرم در سطح اطمینان ۹۰٪ در رودخانه سجاس رود شهر ایجرود مشاهده می‌شود. در ده سال اول و دوم از لحاظ شاخص موران محلی خوشه‌بندی بالا-بالا در رودخانه ایجرود و خوشه‌بندی پایین-پایین در رودخانه زنجان رود مشاهده می‌شود. با توجه به نقشه‌ها می‌توان گفت که احتمال رخداد این نتایج به دلیل نوع زمین‌شناسی و کاربری اراضی منطقه است. هر دو لکه‌ی مشاهده شده در یک ساختار زمین‌شناسی از نوع آهکی قرار دارند بنابراین می‌توان زمین شیمی را عامل تاثیرگذار تشکیل این لکه‌ها دانست که در شکل ۳ مشاهده می‌شود.



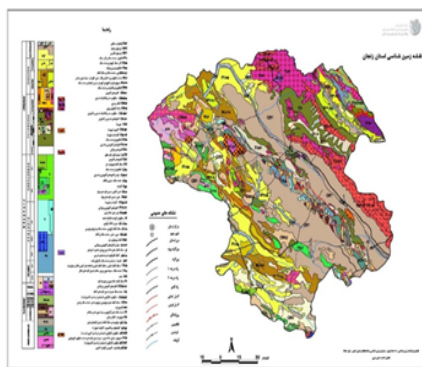
شکل ۱: پراکنش نقاط داغ ترکیب بیکرنات در آب‌های سطحی استان زنجان



شکل ۲: وضعیت خوشه‌بندی بندی ترکیب بیکرنات در آب‌های سطحی استان زنجان

مراجع

- دشتی برمکی، م.، رضایی، م.، صابری نصر، ا. (۱۳۹۳)، ارزیابی شاخص کیفیت آب زیرزمینی (GQI) درآبخوان لنجانان با استفاده سیستم اطلاعات جغرافیایی، نشریه زمین شناسی مهندسی، جلد هشتم، شماره ۲ تابستان.
- فریادی، س.، شاهدهی، ک.، نباتپور، م. (۱۳۹۱)، مطالعه پارامترهای کیفیت آب رودخانه تجن با استفاده از تکنیک‌های آماری چند متغیره، پژوهشنامه مدیریت حوزه آبخیز سال سوم، شماره ۶ پاییز و زمستان.
- بحری، ع.، خسروی، ی.، (۱۳۹۶)، مروری بر آمار فضایی و کاربردهای آن در محیط زیست، چهارمین کنفرانس بین‌المللی برنامه‌ریزی و مدیریت محیط زیست، خرداد ماه.



شکل ۳: زمین شناسی استان زنجان

Ferchichi, H., Hamouda, M. B., Farhat, B., and Mammou, A. B. (2018), Assessment of Groundwater Salinity Using GIS and Multivariate Statistics in a Coastal Mediterranean Aquifer, *International Journal of Environmental Science and Technology*, 1-20.

Li, P., Tian, R., and Liu, R. (2018), Solute Geochemistry and Multivariate Analysis of Water Quality in the Guohua Phosphorite Mine, Guizhou Province, China, *Exposure and Health*, 1-4.

Mohammadjani, E., and Yazdani, N. (2014), Analysis of the Water Crisis and its Management Requirements, *Journal Ravand*, 21, 117-144.



تحلیل زمانی و فضایی- زمانی داده‌های ریسک سلامت در تصادفات جاده‌ای استان قم

حسین والی^۱، امید کریمی، فاطمه حسینی
گروه آمار، دانشگاه سمنان

چکیده: تلفات جاده‌ای یکی از عوامل مهم مرگ غیر طبیعی در ایران و کشورهای در حال توسعه می‌باشد. جلوگیری از بروز تلفات جاده‌ای برای کاهش آمار بالای کشته‌ها امری ضروری به نظر می‌رسد. تعداد مرگ و میر، تعداد مصدومین و تعداد تصادفات جاده‌ای از مولفه‌های مهم ریسک سلامت در تصادفات جاده‌ای به شمار می‌رود. هدف از این مقاله مدل‌سازی نرخ مرگ و میر و مجروحین در محورهای برون شهری استان قم از سال ۱۳۸۸ الی ۱۳۹۵ بر اساس مدل‌های آماری مناسب است. داده‌ها برحسب مکان و زمان وقوع تصادف جمع‌آوری شده‌اند و بعد از پالایش و دسته‌بندی به صورت چهار محور اصلی تردد استان قم (محور شمال - محور جنوب - محور شرق - محور غرب) در هشت سال به‌طور ماهیانه خلاصه شده‌اند. با استفاده از مدل‌های خطی و خطی تعمیم یافته، متغیرهای کمکی تاثیرگذار را به‌وسیله یک تحلیل اکتشافی روی داده‌ها شناسایی کردیم. با توجه به تاثیر زمان و مکان در داده‌ها از مدل‌های سری زمانی چندمتغیره و مدل‌های فضایی- زمانی، برای تحلیل و بررسی آنها استفاده می‌کنیم. در نهایت مدل مناسب برای تحلیل و پیش‌بینی‌های آتی معرفی می‌شود.

واژه‌های کلیدی: داده‌های ریسک سلامت تصادفات استان قم، سری زمانی چندمتغیره، مدل‌های فضایی- زمانی.
کد موضوع بندی ریاضی (۲۰۱۰): 62H11, 62M10, 62M30.

۱ مقدمه

پرداختن به داده‌های نرخ مرگ و میر و مصدومین در تصادفات کشور دارای اهمیت زیادی است. از آنجایی که براساس برآوردهای آماری صورت پذیرفته در طول دو دهه اخیر بیش از ۴۶۰ هزار نفر از هموطنان در کشور جان خود را بر اثر سوانح رانندگی از دست داده‌اند و در عین حال در این ۲۰ سال چیزی نزدیک به ۴/۵ میلیون نفر نیز در جریان تصادفات رانندگی زخمی شده‌اند (والی، ۱۳۹۲). در این مقاله به تحلیل داده‌های ریسک سلامت در تصادفات جاده‌ای استان قم طی سال‌های ۱۳۸۸ الی ۱۳۹۵ می‌پردازیم. مولفه‌های اصلی و موجود برای بررسی ریسک سلامت تصادفات جاده‌ای استان شامل،

^۱ نام و ایمیل ارائه دهنده مقاله: حسین والی، h.vali.4715@gmail.com

تعداد تصادفات، تعداد مجروحین، تعداد مرگ و میرها و نوع وسیله نقلیه (سواری، باربری، اتوبوس و موتورسیکلت) برحسب چهارمحور (شمال، جنوب، شرق، غرب) به صورت ماهیانه است. با توجه به تاثیر زمان در میزان تلفات در تصادفات و حتی تعداد تصادفات، لزوم استفاده از مدل‌هایی که اثر زمان را در نظر بگیرند از اهمیت بالایی برخوردار است. از این رو می‌توان به مدل‌های سری زمانی اشاره کرد که برای تحلیل داده‌هایی به کار می‌روند که برحسب زمان مرتب شده باشند. ابتدا مدل‌های سری زمانی را به تفکیک محورها را مورد مطالعه قرار می‌دهیم، سپس همه محورها را با هم در قالب مدل‌های سری زمانی چندمتغیره (نیرومند، ۱۳۹۱) به کار می‌گیریم. موارد زیادی وجود دارند که مشاهدات مستقل نبوده و برحسب موقعیت قرارگرفتن خود در فضای مورد بررسی به یکدیگر وابسته هستند. اگر این وابستگی تابعی از فاصله بین موقعیت‌های مشاهدات باشد، به گونه‌ای که مشاهدات نزدیک به هم وابسته‌تر و مشاهدات دورتر از هم وابستگی کمتری داشته باشند. اینگونه مشاهدات داده‌های فضایی^۱ نامیده می‌شوند. به طور معمول داده‌های فضایی علاوه بر مقادیر مربوط به متغیر یا متغیرهای مورد بررسی، اطلاع مربوط به مکان مشاهده را نیز شامل می‌شوند. بنابراین منطقی به نظر می‌رسد که مشاهده مربوط به یک مکان به نوعی به مشاهدات سایر مکان‌ها وابسته باشد. داده‌های زمانی نیز دارای چنین خصوصیتی هستند. با این تفاوت که زمان یک جریان یک طرفه و همواره روبه جلو حرکت می‌کند، یعنی از گذشته شروع و از طریق حال به آینده جریان دارد، در حالی که در داده‌های فضایی همبستگی در همه جهات منظور می‌شود. در داده‌های زمانی به دلیل جریان تک بعدی زمان، مشاهدات به طور متوالی به ترتیب زمان رخداد جمع‌آوری و در نتیجه وابستگی به صورت پیاپی بروز می‌کند. اما در داده‌های فضایی، مکان مشاهدات در فضای چند بعدی واقع شده و ممکن است به طور منظم و متوالی روی یک خط قرار نداشته باشند. در این صورت مفهوم ترتیب موقعیت مشاهدات برای داده‌های فضایی معنایی ندارد. بدیهی است که تحلیل آماری داده‌های فضایی با روش‌های آماری معمولی مقدور نیست، زیرا شرط اساسی که همان استقلال داده‌ها است، تحقق پیدا نمی‌کند. از این رو شاخه آمار فضایی برای تحلیل اینگونه مشاهدات شکل گرفته و با فراهم آمدن فنون مختلف درحال توسعه است. علم بشری در طول رشد و شکوفایی خود قدم‌های موثری در این رابطه برداشته است و در این فرآیند، مدل‌سازی را به عنوان یکی از مهمترین راهکارها برای غلبه بر نادانسته‌های خود انتخاب کرده است. یکی از مهمترین روش‌های مدل‌سازی، زمین آمار است که شامل مجموعه‌ای از ابزارهای تکنیکی برای حل مسائل محیطی و ارزیابی‌های زمانی و مکانی پدیده‌های محیطی-فیزیکی مثل عوامل شکل‌گیری و وقوع تصادفات جاده‌ای می‌باشد. برای حل اینگونه مسائل، هر چند می‌توان از روش‌های آمار کلاسیک استفاده کرد، اما مزیت زمین آمار نسبت به آمار کلاسیک در توجه به موقعیت، آرایش و همبستگی بین مشاهدات دقت کافی ندارند. این در حالی است که زمین آمار توجه ویژه‌ای به همبستگی و ساختار مکانی داده‌ها دارد (محمدزاده، ۱۳۹۸؛ کرسی، ۱۳۹۳).

در ادامه بیان می‌کنیم که بعد زمان و مکان داده‌ها به وسیله مدل‌های فضایی-زمانی و سری زمانی چندمتغیره (برحسب محورها) چگونه مورد مطالعه قرار می‌گیرد. در بخش بعدی مدل سری زمانی چندمتغیره را ارائه می‌کنیم. سپس در بخش ۳ یک تحلیل اکتشافی از داده‌ها بیان می‌شود، در بخش ۴ مدل فضایی-زمانی را روی داده‌ها به کار می‌گیریم و در نهایت یک بحث و نتیجه‌گیری ارائه می‌شود.

۲ مدل‌های سری زمانی چندمتغیره

در بسیاری از مطالعات تجربی داده‌های سری زمانی مشاهدات مربوط به چندین متغیر را شامل می‌شود. برای مثال، در مطالعه چگونگی فروش، متغیرها، می‌توانند حجم فروش، قیمت‌ها، توان فروش و هزینه تبلیغات را در بر داشته باشد. در الگوی تابع تبدیل رابطه ویژه بین متغیرها، یعنی یک رابطه ورودی-خروجی بین یک متغیر خروجی و یک یا چند متغیر ورودی را مطالعه

^۱spatial data

می‌کنیم. با وجود این در بسیاری از زمینه‌های کاربردی یک الگوی تابع تبدیل از نوع خروجی، ممکن است مناسب نباشد. در این بخش یک رده کلی الگوی سری زمانی برداری را برای بیان روابط میان چندمتغیر سری زمانی معرفی می‌کنیم.

۱.۲ توابع ماتریس کوواریانس و همبستگی

فرض کنید $\mathbf{X}_t = [X_{1,t}, X_{2,t}, \dots, X_{m,t}]'$ و $t = 0, \pm 1, \pm 2, \dots$ فرآیند برداری با مقدار حقیقی تواما ایستای m - بعدی باشد که میانگین $E(X_{i,t}) = \mu_i$ برای هر $i = 1, 2, \dots, m$ ثابت و کوواریانس متقابل بین $X_{j,s}$ و $X_{i,t}$ برای تمام $i = 1, 2, \dots, m$ و $j = 1, 2, \dots, m$ و $t = 0, \pm 1, \pm 2, \dots$ توابع فقط از اختلاف زمان $(s - t)$ است. بنابراین بردار میانگین به صورت

$$E(\mathbf{X}_t) = \boldsymbol{\mu} = \begin{bmatrix} \mu_1 \\ \mu_2 \\ \vdots \\ \mu_m \end{bmatrix}, \quad (1.2)$$

و ماتریس کوواریانس به صورت

$$\begin{aligned} \Gamma(k) &= \text{Cov}(X_t, X_{t+k}) = E((X_t - \boldsymbol{\mu})(X_{t+k} - \boldsymbol{\mu})'), \\ &= E \begin{bmatrix} X_{1,t} - \mu_1 \\ X_{2,t} - \mu_2 \\ \vdots \\ X_{m,t} - \mu_m \end{bmatrix} [X_{1,t+k} - \mu_1, X_{2,t+k} - \mu_2, \dots, X_{m,t+k} - \mu_m], \\ &= \begin{bmatrix} \gamma_{11}(k) & \gamma_{12}(k) & \dots & \gamma_{1m}(k) \\ \gamma_{21}(k) & \gamma_{22}(k) & \dots & \gamma_{2m}(k) \\ \vdots & \vdots & \dots & \vdots \\ \gamma_{m1}(k) & \gamma_{m2}(k) & \dots & \gamma_{mm}(k) \end{bmatrix} = \text{Cov}(X_{t-k}, X_t), \end{aligned} \quad (2.2)$$

است که به ازای $k = 0, \pm 1, \pm 2, \dots$ و $i = 1, 2, \dots, m$ داریم:

$$\begin{aligned} \gamma_{ij}(k) &= E[(X_{i,t} - \mu_i)(X_{j,t+k} - \mu_j)], \\ &= E[(X_{i,t-k} - \mu_i)(X_{j,t} - \mu_j)]. \end{aligned}$$

به عنوان تابعی از k ، $\Gamma(k)$ تابع ماتریس کوواریانس بردار فرآیند X_t نامیده می‌شود. برای $i = j$ ، تابع کوواریانس برای مؤلفه i ام فرآیند یعنی $X_{i,t}$ و $X_{j,t}$ بوده و برای $i \neq j$ ، تابع کوواریانس متقابل بین $X_{j,t}$ و $X_{i,t}$ است. به سهولت می‌توان دید که $\Gamma(0)$ ماتریس واریانس-کوواریانس فرآیند است. باید توجه کنیم که در فرآیند تواما ایستا نتیجه می‌شود که هر مؤلفه یک متغیر فرآیند ایستاست. با وجود این بردار فرآیندهای ایستای یک متغیری، الزاماً یک فرآیند تواما ایستای برداری نیست.

تابع ماتریس همبستگی برای فرآیند برداری به صورت

$$\rho(k) = D^{-\frac{1}{2}} \Gamma(k) D^{-\frac{1}{2}} = [\rho_{ij}(k)], \quad (3.2)$$

برای $i = 1, \dots, m$ و $j = 1, \dots, m$ تعریف می‌شود، که D ماتریس قطری است که در آن عنصر قطری i ام واریانس فرآیند i ام است، یعنی:

$$D = \text{diag} [\gamma_{11}(0), \gamma_{22}(0), \dots, \gamma_{mm}(0)].$$

واضح است که عنصر قطری i ام $\rho(k)$ تابع خودهمبستگی مولفه i ام سری یعنی $X_{i,t}$ است. در حالی که عنصر غیر قطری (i, j) ، یعنی:

$$\rho_{ij}(k) = \frac{\gamma_{ij}(k)}{[\gamma_{ij}(0)\gamma_{ij}(0)]^{\frac{1}{2}}}, \quad (4.2)$$

است، که تابع خودهمبستگی متقابل بین $X_{j,t}$ و $X_{i,t}$ را نشان می‌دهد.

۲.۲ الگوهای اتورگرسیو میانگین متحرک برداری نایستا

در تحلیل سری‌های زمانی، مشاهده سری‌هایی که رفتاری نایستا را نشان می‌دهند، بسیار متداول است. یک راه مفید تبدیل سری‌های نایستا به ایستا، تفاضلی کردن است. برای مثال، در سری‌های زمانی یک متغیری یک سری نایستای X_t به یک سری ایستا $(1-B)^d X_t$ ، $d > 0$ تبدیل می‌شود، به‌طوری که می‌توان نوشت:

$$\phi_p(B)(1-B)^d X_t = \theta_q(B)Z_t, \quad (5.2)$$

که در آن $\phi_p(B)$ یک عملگر اتورگرسیو^۲ (AR) ایستاست. چنین به نظر می‌رسد که تعمیم (۵.۲) فرآیند برداری:

$$\Phi_p(B)(I - IB)^d \mathbf{X}_t = \Theta_q(B)\mathbf{Z}_t, \quad (6.2)$$

است. حال از این تعمیم، نتیجه می‌شود که تمام مؤلفه‌های سری به دفعات مساوی تفاضلی شده است. روشن است که این محدودیت لازم نبوده و ناخواسته است. برای انعطاف پذیری بیشتر، فرض می‌کنیم که $\mathbf{X}_t$ ایستا نباشد ولی آن را با به کار بردن عملگر تفاضلی زیر می‌توان به ایستا تبدیل نمود:

$$D(B)\mathbf{X}_t, \quad (7.2)$$

که در آن

$$D(B) = \begin{bmatrix} (1-B)^{d_1} & \dots & \dots & \dots \\ \dots & (1-B)^{d_2} & \dots & \dots \\ \vdots & \vdots & \ddots & \vdots \\ \dots & \dots & \dots & (1-B)^{d_m} \end{bmatrix}, \quad (8.2)$$

و $(d_1, \dots, d_m)$ مجموعه‌ای از اعداد صحیح نامنفی است. لذا الگوی^۳ (ARMA) برداری نایستا به صورت

$$\Phi_p D(B)\mathbf{X}_t = \Theta_q(B)\mathbf{Z}_t, \quad (9.2)$$

²Autoregressive

³Autoregressive Moving Average

را داریم، که ریشه‌های صفر $|\Phi_p(B)|$ و $|\Theta(B)|$ خارج دایره واحدند. تفاضلی کردن سری‌های زمانی برداری بسیار پیچیده است و باید با دقت انجام شود. زیاد تفاضلی کردن هر مؤلفه سری، ممکن است منتهی به پیچیدگی‌هایی در برازش الگو شود. ثابت می‌شود که به طور یکسان تفاضلی کردن پیشنهاد شده برای یک فرآیند برداری، ممکن است نتیجه‌اش یک نمایش وارون‌پذیر نباشد. زیرا یک اریبی بزرگ در برآورد پارامترها، ممکن است از یک الگوی وارون‌پذیر ناشی شود، که از زیاد تفاضلی کردن باید پرهیز شود (نیرومند، ۱۳۹۱). از این رو، تعمیم زیر برای الگوی ARMA برداری:

$$\Phi_p(B)\mathbf{X}_t = \Theta_q(B)\mathbf{Z}_t, \quad (10.2)$$

پیشنهاد شده است (نیرومند، ۱۳۹۱)، که صفرهای چندجمله‌ای دترمینانی $\Phi_p(B)$ رو یا خارج دایره واحدند. در ادامه مدل‌های ARMA (۹.۲) را در حالت نایستای بررسی می‌کنیم.

۳.۲ مدل اتو رگرسیو میانگین متحرک تلفیق شده

مدل اتو رگرسیو میانگین متحرک تلفیق شده ^۴ (ARIMA) که فرآیند ایستای حاصل از یک سری نایستای همگن تفاضلی است. به‌طوری که سری تفاضلی شده $(1-B)^d X_t$ از یک فرآیند ایستای $ARMA(p, q)$ پیروی می‌کند. بنابراین داریم:

$$\phi_p(B)(1-B)^d X_t = \theta_0 + \theta_q(B)Z(t), \quad (11.2)$$

که در آن عملگر AR ایستای $\phi_p(B) = (1 - \phi_1 B - \dots - \phi_p B^p)$ و عملگر MA وارون‌پذیر

$$\theta_q(B) = (1 - \theta_1 B - \dots - \theta_q B^q),$$

دارای ریشه مشترکی نیستند. الگوی نایستای همگن حاصل در (۱۱.۲) را یک الگوی اتو رگرسیو تلفیق شده با مرتبه (p, d, q) می‌نامند که به صورت $ARIMA(p, d, q)$ نشان داده می‌شود. وقتی $p = *$ الگوی $ARIMA(p, d, q)$ را یک الگوی میانگین متحرک تلفیق شده از مرتبه (d, q) می‌نامند که به صورت الگوی $IMA(d, q)$ نشان داده می‌شود. $q = *$ فرآیند یک الگوی برداری $AR(p)$ به صورت

$$\mathbf{X}_t = \Phi_1 \mathbf{X}_{t-1} - \dots - \Phi_p \mathbf{X}_{t-p} + \mathbf{Z}_t, \quad (12.2)$$

می‌شود. فرآیند ایستاست هرگاه ریشه‌های چندجمله‌ای دترمینانی $\Phi_q(B)$ خارج دایره واحد باشند. برای اطلاعات بیشتر در مورد مدل‌های سری‌زمانی چندمتغیره (نحوه برازش مدل و پیشگویی) به نیرومند (۱۳۹۱) مراجعه شود. در ادامه به تحلیل اکتشافی داده‌ها می‌پردازیم و این مدل‌ها را روی آنها به کار می‌گیریم.

۳ تحلیل اکتشافی داده‌ها

در این بخش با استفاده از مدل‌های خطی و خطی تعمیم‌یافته به بررسی هر یک از متغیرها روی متغیر پاسخ می‌پردازیم. با این رویکرد که متغیرهای تاثیرگذار شناسایی شوند. داده‌ها برحسب مکان و زمان وقوع تصادف براساس اطلاعات حاصل از آمارهای رسمی و ثبتی پلیس راه و پزشکی قانونی استان قم جمع‌آوری شده‌اند. این داده‌ها شامل متغیرهای سال (year)، ماه (months)، موقعیت محورها (location)، تعداد تصادفات (Accidents)، تعداد مصدومین (Injured)، تعداد مرگ و میر (Fatalities)، تعداد وسیله نقلیه سواری‌ها در تصادف (Riding)، تعداد موتور سیکلت‌ها در تصادف (motorcycle)،

⁴ Autoregressive integrated Moving Average

عابر پیاده (pedestrian)، وسیله نقلیه باری (load)، اتوبوس (Bus)، عرض جغرافیایی محور تصادف (lat)، طول جغرافیایی محور تصادف (long) و نرخ مرگ و میر به تعداد تصادف (Rate) است.

در حالت اول متغیر تعداد مرگ و میرها را به عنوان پاسخ و بقیه متغیرها به عنوان متغیر کمکی، مستقل یا کمی در نظر گرفته ایم به این صورت است که متغیرهای سال، ماه و موقعیت محور به عنوان متغیرهای کیفی طبقه بندی شده را وارد مدل کرده ایم و متغیرهای تعداد تصادفات، تعداد مصدومان و ... به عنوان متغیرهای کمی در مدل لحاظ کرده ایم. خلاصه برازش مدل این بود که متغیرهای وسیله نقلیه سواری، باربری، اتوبوس و عابر پیاده در سطح ۵٪ معنی دار هستند. البته با ملاک اطلاع آکائیک (AIC) مدل نهایی به صورت

$$\log E(Fatalities) \sim Injured + Riding + motorcycle + pedestrian + load + Bus + \varepsilon_t$$

حاصل شد. هر چند شرایط مدل با بررسی خطاها برقرار نشد و لذا مدل فوق ناکارآمد می باشد. در حالت دوم متغیر نرخ مرگ و میر بر تعداد تصادفات را به عنوان متغیر پاسخ در نظر گرفتیم و همانند مدل قبلی متغیرهای سال، ماه و موقعیت محور به عنوان متغیر کیفی وارد مدل کردیم. خلاصه نتایج برازش مدل به این صورت بود که متغیرهای محور، ماه، انواع وسیله نقلیه و عابر پیاده در سطح ۵٪ معنی دار هستند. البته در این حالت نیز شرایط مدل با بررسی باقیمانده ها مورد تایید نیست. بنابراین با اطلاعاتی که کسب کردیم به این نتیجه رسیدیم که محورها را به تفکیک مورد بررسی قرار دهیم، لذا در حالت سوم به این مورد پرداختیم که تحلیل سری زمانی به تفکیک محورها با در نظر گرفتن متغیر تعداد مرگ و میر به عنوان متغیر پاسخ و متغیرهای کمکی تعداد تصادفات، تعداد مصدومین، انواع وسیله های نقلیه انجام دهیم. در واقع یک مدل خطی تعمیم یافته با خطاهای سری زمانی $ARMA$ به داده ها برازش دهیم. مدل کلی (M Full) بدون تفکیک محورها به صورت

$$\log E(Fatalities) \sim Accident + Injured + Riding + motorcycle + pedestrian + load + Bus + \varepsilon_t$$

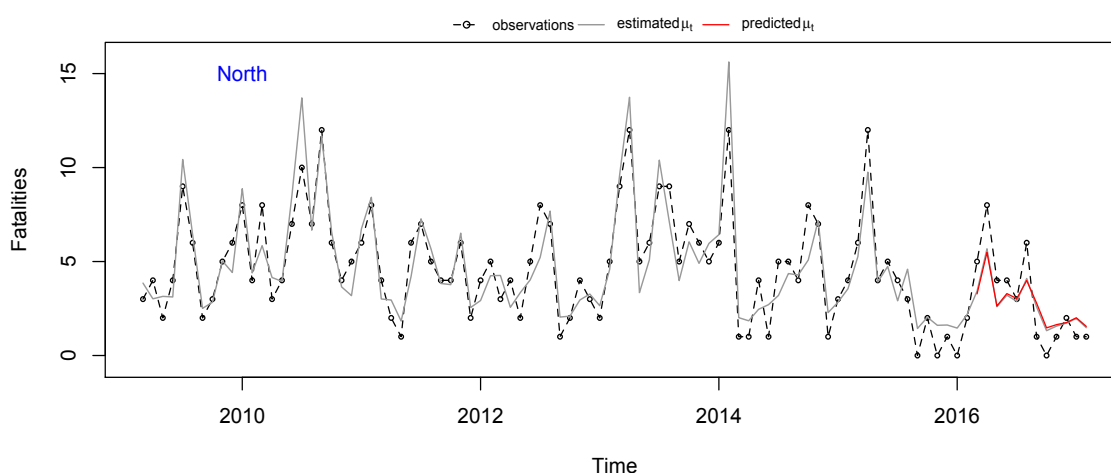
است، که در آن خطای ε_t را به صورت یک مدل سری زمانی $ARMA$ در نظر می گیریم. با توجه به اینکه داده ها در چهار محور (شمال، جنوب، شرق، غرب) استان ثبت شده است برای هر محور مدل به طور جداگانه برازش می شود. با نمادهای، مدل محور شمال (M North)، مدل محور جنوب (M South)، مدل محور غرب (M West) و مدل محور شرق (M East) نمایش می دهیم. نتایج در جدول ۱ خلاصه شده است: در جدول ۱ مقادیر برآورد ضرایب مدل (Est) و p - مقدار آنها

جدول ۱: نتایج برآورد پارامترهای مدل و p - مقدار برای مدل کلی و به تفکیک چهار محور جاده ای قم

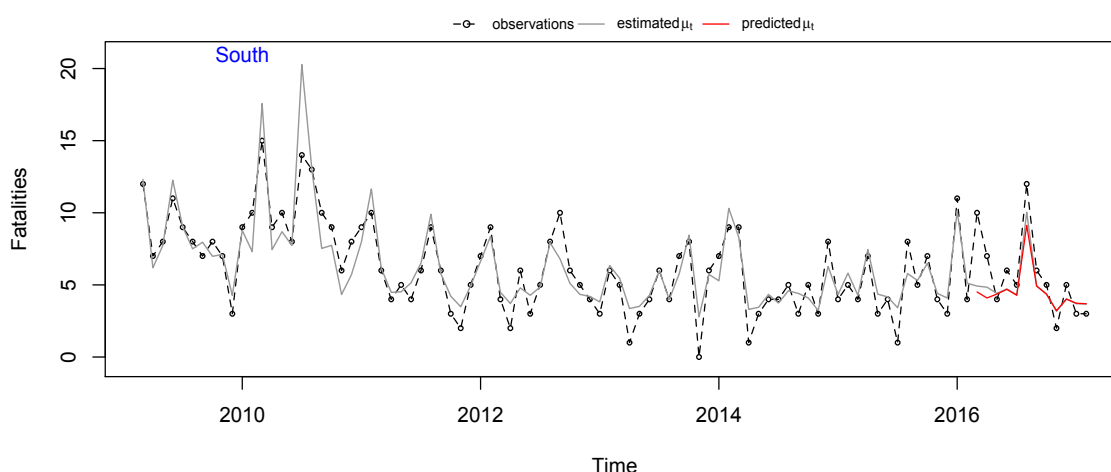
	M . Full		M . North		M . South		M . West		M . East	
	<i>Est</i>	<i>p</i>	<i>Est</i>	<i>p</i>	<i>Est</i>	<i>p</i>	<i>Est</i>	<i>p</i>	<i>Est</i>	<i>p</i>
<i>Intercept</i>	۰/۶۸۹۱	۰/۰۰	۰/۱۰۱۷	۰/۷۹	۰/۷۸۷۹	۰/۰۲	۰/۳۶۶	۰/۲۰	۰/۱۰۷	۰/۷۹
<i>Accident</i>	-۰/۰۰۱۵	۰/۵۳	۰/۰۰۴۶	۰/۴۶	۰/۰۰۴۸	۰/۴۹	-۰/۰۰۱	۰/۸۶	۰/۰۰۲۹	۰/۷۷
<i>Injured</i>	۰/۰۰۸۵	۰/۱۲	۰/۱۱۸	۰/۲۵	۰/۰۰۰۸	۰/۹۴	۰/۰۱۱	۰/۳۵	۰/۰۰۹۱	۰/۵۰
<i>Riding</i>	۰/۱۸۵۶	۰/۰۰	۰/۲۲۶۶	۰/۰۰	۰/۱۴۱۶	۰/۰۰	۰/۲۶۳	۰/۰۰۷	۰/۲۹۶۲	۰/۰۰
<i>Motor</i>	۰/۱۳۳۲	۰/۱۵۸	۰/۱۴۳۹	۰/۴۶	۰/۰۵۳۸	۰/۷۰	۱/۱۸۷	۰/۰۲	۰/۰۹۱۴	۰/۷۱
<i>Pedes</i>	۰/۲۸۹۴	۰/۰۰۶	۰/۰۴۶۹	۰/۸۳	۰/۱۵۴۹	۰/۴۲	۰/۲۱۲	۰/۶۱	۰/۰۱	۰/۲۱۶
<i>Load</i>	۰/۱۷۰۴	۰/۰۰	۰/۱۳۷۹	۰/۱۲	۰/۱۷۶۷	۰/۰۱	۰/۸۴۱	۰/۳۶	۰/۰۱۵۴	۰/۰۱۵
<i>Bus</i>	۰/۲۳۷۲	۰/۰۰۲	۰/۰۱۸۲	۰/۹۳	۰/۰۷۹۶	۰/۵۹	۰/۱۷۱	۰/۲۹	۰/۳۲۳۶	۰/۱۶
<i>MSE</i>	۲/۸۰۱۳۱۱		۱/۵۶۲۸		۲/۰۸۴		۱/۲۶۶۵		۱/۲۲۲۰	
<i>AIC</i>	۱۴۵۵/۹۴۸		۳۰۶/۴۶		۳۹۴/۱۴۱		۳۵۹/۰۰۳۵		۳۴۱/۵۸	

(p) برای مدل کلی و به تفکیک چهار محور ارائه شده است و همچنین مقدار MSE و AIC برای مقایسه بین مدل ها نیز

محاسبه شده‌اند. می‌توان نتیجه گرفت که وقتی مدل کلی را بدون توجه به چهار محور اجرا کردیم هم مقدار MSE ^۵ و هم AIC ^۶ آن نسبت به حالت تفکیک شده خیلی بیشتر است یعنی اینکه بدون توجه اطلاع در مورد چهار محور دقت مدل پایین می‌آید. با توجه به نتایج برازش چهار محور در جدول ۱، یک نتیجه‌گیری واحد در مورد متغیرهای مورد نظر در مدل نمی‌توان گرفت. اما براساس مقادیر MSE و AIC می‌توان گفت که محور شرق MSE کمتر و محور شمال AIC کمتری دارند. شکل‌های ۱ تا ۴ نمودار سری زمانی مشاهدات، مقادیر برازش شده و مقادیر پیش‌بین برای سال ۲۰۱۶ را نشان می‌دهند. که بیانگر این است چگونه مدل به مشاهدات به خوبی برازش شده‌اند و همچنین پیش‌بینی آینده نیز با واقعیت تا حدودی هم‌خوانی دارد.

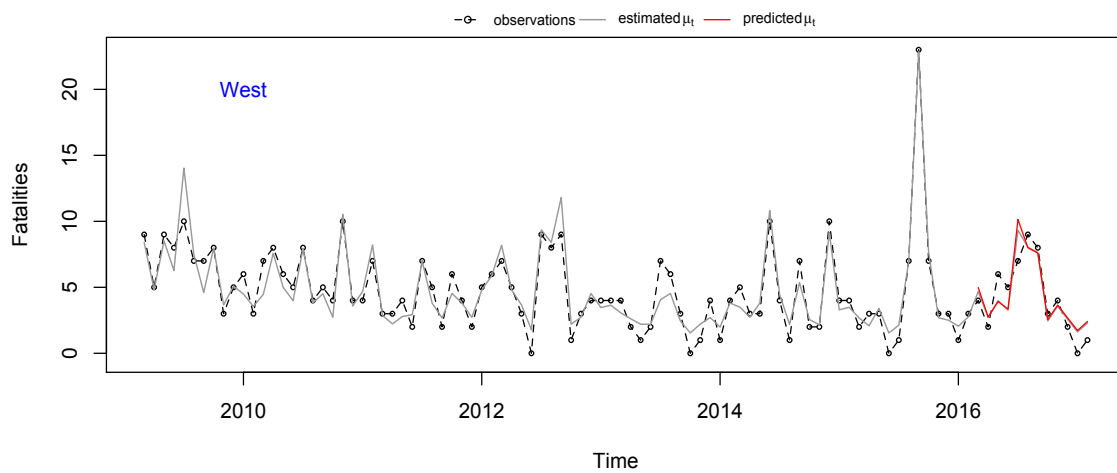


شکل ۱: سری زمانی تعداد مرگ و میرهای تصادفات (شمال): مقادیر برازش شده (خط چین)، مقادیر پیش‌بینی آینده (خط قرمز)

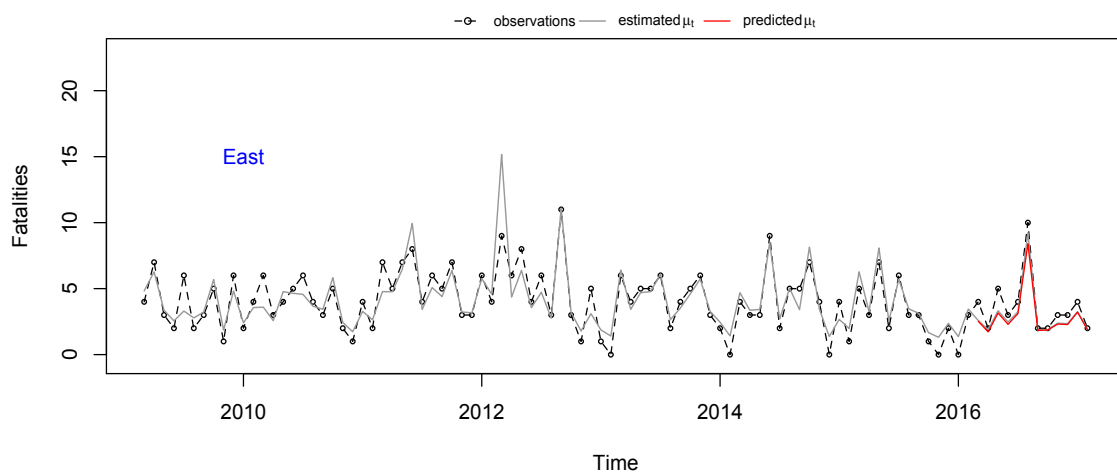


شکل ۲: سری زمانی تعداد مرگ و میرهای تصادفات (جنوب): مقادیر برازش شده (خط چین) و مقادیر پیش‌بینی آینده (خط قرمز)

^۵ Akaike information criterion



شکل ۳: سری زمانی تعداد مرگ و میرهای تصادفات (غرب): مقادیر برازش شده (خط چین) و مقادیر پیش‌بینی آینده (خط قرمز)



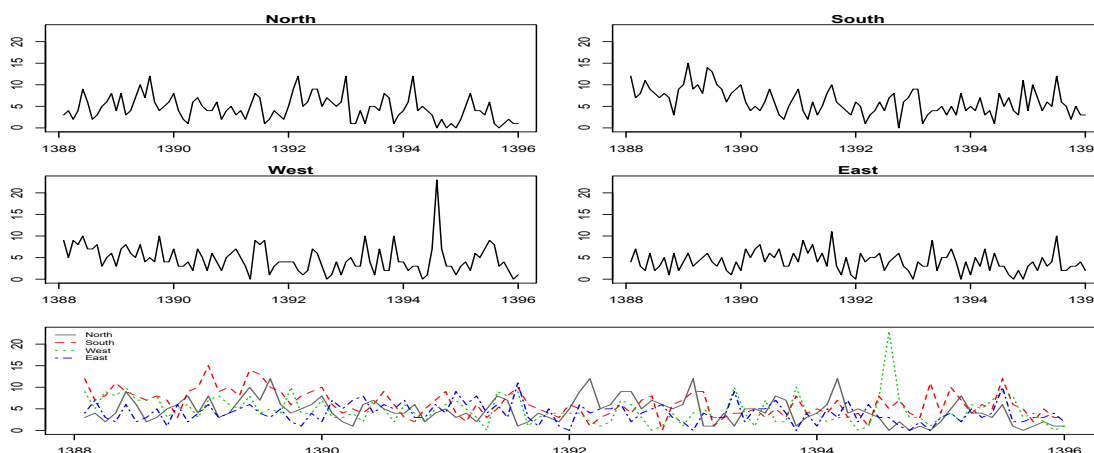
شکل ۴: سری زمانی تعداد مرگ و میرهای تصادفات (شرق): مقادیر برازش شده (خط چین) و مقادیر پیش‌بینی آینده (خط قرمز)

۱.۳ تحلیل سری زمانی چند متغیره داده‌های مرگ و میر تصادفات جاده‌ای استان قم

در این بخش به تحلیل سری زمانی چند متغیره روی داده‌های تعداد مرگ و میر تصادفات جاده‌ای می‌پردازیم. در اینجا هر محور را به عنوان یک متغیر پاسخ در نظر گرفته‌ایم، بنابراین یک متغیر پاسخ (چهار متغیره) برداری به صورت

$$Z_t = \begin{bmatrix} z_{1t} \\ z_{2t} \\ z_{3t} \\ z_{4t} \end{bmatrix}, \quad t = 1, 2, \dots, 12 * 8$$

داریم. نمودار سری زمانی تعداد مرگ و میرها برای چهار محور در شکل ۵ ارائه شده است. همانطور که از نمودار برداشت می‌شود، تعداد مرگ و میرها در محور جنوب روند کاهشی دارد و در محور شرق تغییرات کمتر می‌باشد. تعداد مرگ و میرها



شکل ۵: نمودار سری زمانی تعداد مرگ و میرهای تصادفات در چهار محور شمال، جنوب، غرب و شرق

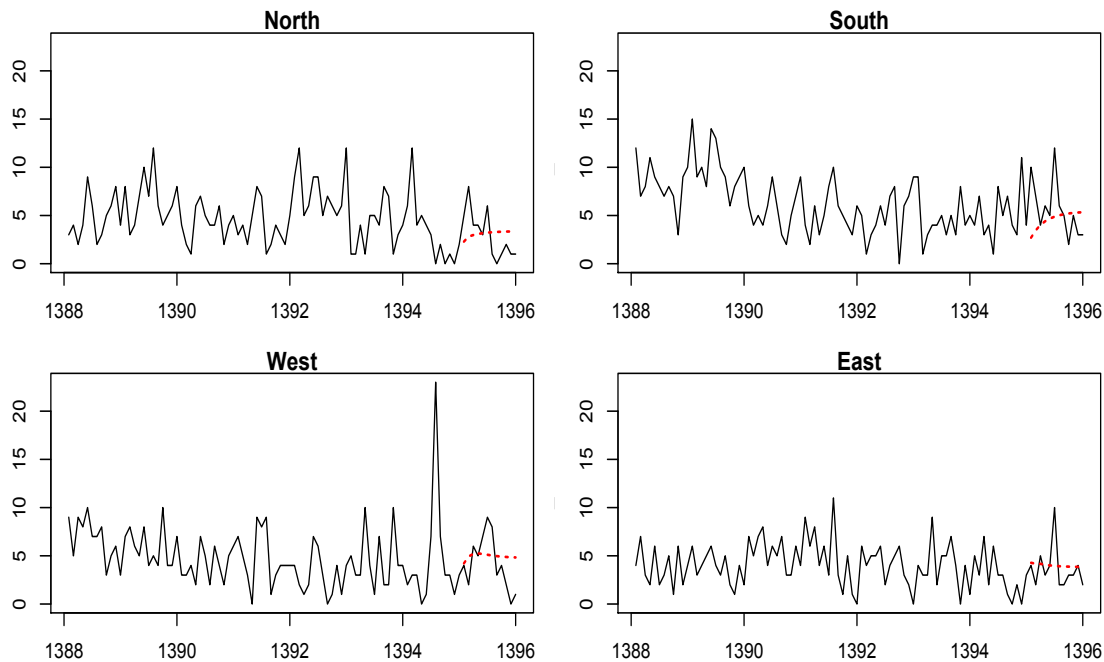
در محور جنوب از سال ۱۳۸۸ تا سال ۱۳۹۶ یک روند کاهشی مشاهده می‌شود و در محور شرق تغییرات کمتری نسبت به محورهای دیگر به نظر می‌رسد. بیشترین کشته‌ها در تصادفات مربوط به محور غرب استان در سال ۹۴ مشاهده می‌شود. مدل‌های سری زمانی چندمتغیره^۷ $VARMA(p, q)$ با مرتبه‌های متفاوت روی داده‌ها بوسیله نرم‌افزار R با بسته MTS پیاده‌سازی کردیم. نتایج حاصل با در نظر گرفتن معیار $RMSE$ و AIC انتخاب مدل‌ها در جدول ۲ خلاصه شده‌اند. با

جدول ۲: مقادیر $RMSE$ و AIC برای مدل $VARMA$ با مرتبه‌های مختلف

Model	AIC	RMSE
$VARMA(0, 1)$	۷/۹۸۵۶	۲/۶۵۸۸
$VARMA(1, 0)$	۷/۹۵۶۵	۲/۶۵۰۳
$VARMA(1, 1)$	۷/۸۶۱۸	۲/۵۳۸۴
$VARMA(2, 1)$	۸/۱۵۸۲	۲/۵۱۸۴
$VARMA(1, 2)$	۸/۱۰۱۸	۲/۵۱۵۷
$VARMA(2, 2)$	۸/۰۳۹۹	۲/۳۷۷۰

توجه به مقادیر $RMSE$ و AIC مدل $VARMA(1, 1)$ را انتخاب کردیم چون هم AIC کمتری دارد و هم خلاصه‌تر از مدل $VARMA(2, 2)$ از نظر تعداد پارامتر می‌باشد. شکل ۶ نمودار سری زمانی داده‌ها و مقادیر پیش‌بینی برای مدل $VARMA(1, 1)$ در ۱۲ ماه سال ۱۳۹۵ را نمایش می‌دهد. با توجه به شکل می‌توان نتیجه گرفت که پیش‌بینی به خوبی صورت نگرفته است. برای پیش‌بینی مناسب، بهتر است متغیرهای کمکی تعداد تصادفات، تعداد مصدومین و وسیله نقلیه سواری در مدل لحاظ شود. برای اینکار ابتدا حذف روند به وسیله یک مدل رگرسیونی چند متغیره خطی انجام داده‌ایم سپس مدل $VARMA(p, q)$ را برای داده‌های بدون روند برازش کرده‌ایم و نتایج $RMSE$ و AIC انتخاب مدل‌ها در جدول ۳ خلاصه شده‌اند. لازم به ذکر است که مدل‌های $VARMA$ با مرتبه همزمان p و q یعنی $VARMA(1, 1)$ ، $VARMA(2, 1)$ و $VARMA(2, 2)$ از نظر محاسباتی همراه با سه متغیر کمکی قابل انجام نبود یعنی برآورد ماکسیمم

⁷Vector Autoregressive integrated Moving Average

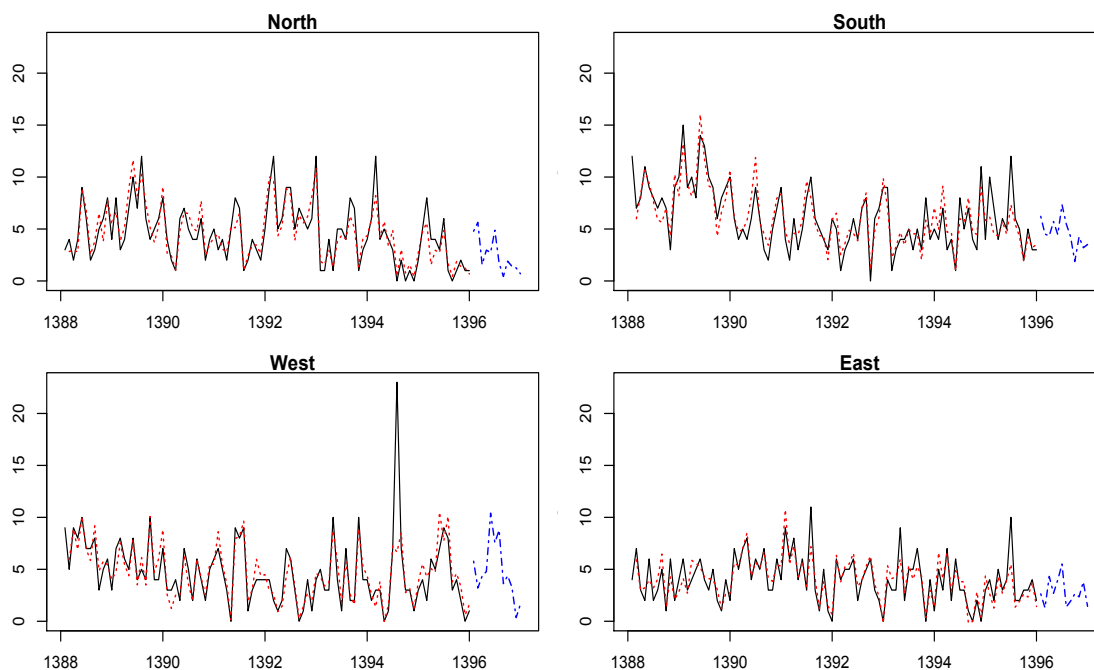


شکل ۶: نمودار سری داده‌های مرگ و میر به همراه مقادیر پیش‌بینی در ۱۲ ماه سال ۱۳۹۶ (خط چین)

جدول ۳: مقادیر AIC و $RMSE$ برای مدل بدون روند $VARMA(p, q)$

Model	AIC	RMSE
$VARMA(0, 1)$	۲/۷۰۰۲	۱/۴۲۶۸
$VARMA(1, 0)$	۲/۷۰۰۱	۱/۴۲۴۷
$VARMA(0, 2)$	۲/۸۹۲۲	۱/۴۰۴۶
$VARMA(2, 0)$	۲/۹۱۱۷	۱/۴۰۵۴

درست‌نمایی پارامترها به روش عددی قابل محاسبه نیست، زیرا تعداد پارامترها خیلی زیاد است. اما با دو متغیر کمکی، تعداد تصادفات و تعداد مصدومان قابل انجام بود ولی مقادیر AIC آنها بیشتر از عدد هفت شد و عملاً قابل رقابت با مدل‌های جدول ۳ نیستند، بنابراین مناسب‌ترین مدل با توجه به نتایج جدول ۳ مدل $VARMA(1, 0)$ است. شکل ۷ نمودار سری زمانی داده‌ها به همراه مقادیر برازش شده سری از مدل $VARMA(1, 0)$ و مقادیر پیش‌بینی برای ۱۲ ماه سال ۱۳۹۶ را نمایش می‌دهد و حاکی از برازش مناسب $VARMA(1, 0)$ بدون روند نسبت به مدل $VARMA(1, 1)$ شکل ۶ است. اکنون که مناسب‌ترین مدل از بین مدل‌ها و روش‌های پیشنهادی مشخص گردید، می‌توان براساس پیش‌بینی برای تعداد مرگ و میر ماه‌های آتی با دقت بیشتری انجام داد البته لازم به ذکر است همه این پیش‌بینی‌ها فقط زمانی قابل اتکا هستند که همان شرایطی که برای سال‌های گذشته بوده است نیز برای سال‌های آتی برقرار باشد. دربخش بعد موقعیت محورهای رابه‌طور مستقیم به‌وسیله مدل‌های فضایی-زمانی در نظر گرفته‌ایم و نتایج را با هم مورد مقایسه قرار داده‌ایم.



شکل ۷: نمودار سری زمانی مشاهدات به همراه مقادیر برازش شده (خط چین) و مقادیر پیش‌بینی سال ۱۳۹۶

۴ مدل‌های فضایی-زمانی

فرض کنید $z_\ell(s_i, t)$ مقادیر مشاهدات پاسخ در موقعیت s_i ، $i = 1, \dots, n$ ، و ماه t ام در سال ℓ ام برای $\ell = 1, \dots, r$ و $t = 1, \dots, T_\ell$ باشد. در اینجا دو اندیس زمان به صورت ماه درون سال مشخص شده است که می‌تواند به صورت روز درون سال باشد. قرار دهید $z_{\ell t} = (z_\ell(s_1, t), \dots, z_\ell(s_n, t))'$ و فرض کنید $x_{\ell j}(s, t)$ متغیر کمکی j ام در ماه t ام سال ℓ ام باشد و قرار دهید $X_{\ell t}$ ماتریس $n \times p$ از متغیرهای که شامل عرض از مبدأ است. **گلفند (۲۰۱۲)** و **گلفند و همکاران (۲۰۰۵)** مدل‌های فضایی-زمانی بیزی را به صورت سلسله‌مراتبی در سه مرحله بیان کرد: ۱- مشاهدات به شرط یک فرآیند تصادفی و پارامترها $(z|u, \theta)$ ، ۲- فرآیند تصادفی به شرط پارامتر $(u|\theta)$ ، ۳- پارامترها (θ) . در مرحله اول مدل به صورت

$$z_{\ell t} = u_{\ell t} + \varepsilon_{\ell t}, \quad \ell = 1, \dots, r, \quad t = 1, \dots, T_\ell,$$

در نظر گرفته می‌شود، که در آن $\varepsilon_{\ell t} = (\varepsilon_\ell(s_1, t), \dots, \varepsilon_\ell(s_n, t))'$ بردار خطاهای مدل با توزیع $N(\cdot, \sigma_\varepsilon^2 I_n)$ است و σ_ε^2 اثر قطعه‌ای یا واریانس خطای تصادفی است. $u_{\ell t} = (u_\ell(s_1, t), \dots, u_\ell(s_n, t))'$ مقادیر واقعی متناظر با $z_\ell(s_i, t)$ که شامل مولفه‌های $u_{\ell t} = X_{\ell t}\beta + \eta_{\ell t}$ است، که در آن $\beta = (\beta_1, \dots, \beta_p)'$ بردار پارامترهای ضرایب رگرسیونی، $\eta_{\ell t} = (\eta_\ell(s_1, t), \dots, \eta_\ell(s_n, t))'$ اثر تصادفی فضایی-زمانی مدل است که دارای توزیع $N(\cdot, \Sigma_n)$ و مستقل از زمان می‌باشد، که در آن $\sigma_\eta^2 = \sigma_\eta^2 S_\eta$ پارامتر واریانس فضایی و S_η ماتریس همبستگی فضایی است. در اینجا از تابع همبستگی ماترن (ماترن، ۱۹۸۶) به صورت

$$K(s_i, s_j; \varphi, \nu) = \frac{1}{2^{\nu-1} \Gamma(\nu)} (\sqrt{\nu} \|s_i - s_j\| \varphi)^\nu K_\nu(\sqrt{\nu} \|s_i - s_j\| \varphi)$$

استفاده شده است، که در آن K_ν تابع بسل اصلاح شده نوع دوم با مرتبه ν و $\|s_i - s_j\|$ فاصله بین موقعیت‌های s_i و s_j است. φ پارامتر دامنه فضایی و ν پارامتر هموارساز میدان تصادفی است، (کرسی، ۱۳۹۳؛ کرسی و ویکل، ۲۰۱۱).

مدل سلسله‌مراتبی ذکر شده را مدل فرآیند گاوسی (GP) مستقل می‌گویند. فرض کنید $\theta = (\beta, \sigma_\varepsilon^2, \sigma_\eta^2, \varphi, \nu)$ بردار پارامترهای مدل، $z_{\ell t}^{obs}$ بردار داده‌های مشاهده شده و $z_{\ell t}^*$ بردار داده‌های گمشده باشد. با در نظر گرفتن توزیع‌های پیشین معمول برای پارامترهای مدل؛ $\pi(\theta)$ ، لگاریتم توزیع پسین توأم پارامترها و داده‌های گمشده برای مدل GP به صورت

$$\begin{aligned} \log \pi(\theta, \mathbf{u}_{\ell t}, \mathbf{z}_{\ell t}^* | \mathbf{z}_{\ell t}^{obs}) \propto & -\frac{N}{2} \log \sigma_\varepsilon^2 - \frac{1}{2\sigma_\varepsilon^2} \sum_{\ell=1}^r \sum_{t=1}^{T_\ell} (\mathbf{z}_{\ell t} - \mathbf{u}_{\ell t})' (\mathbf{z}_{\ell t} - \mathbf{u}_{\ell t}) \\ & - \frac{\sum_{\ell=1}^r T_\ell}{2} \log |\sigma_\eta^2 S_\eta| \\ & - \frac{1}{2\sigma_\eta^2} \sum_{\ell=1}^r \sum_{t=1}^{T_\ell} (\mathbf{u}_{\ell t} - \mathbf{X}_{\ell t} \beta)' S_\eta^{-1} (\mathbf{u}_{\ell t} - \mathbf{X}_{\ell t} \beta) + \log \pi(\theta) \end{aligned}$$

به دست می‌آید. مدل سلسله‌مراتبی اتورگرسیو (AR) به صورت

$$\begin{aligned} \mathbf{z}_{\ell t} &= \mathbf{u}_{\ell t} + \varepsilon_{\ell t}, \\ \mathbf{u}_{\ell t} &= \rho \mathbf{u}_{\ell, t-1} + \mathbf{X}_{\ell t} \beta + \eta_{\ell t}, \end{aligned}$$

معرفی می‌شوند، که در آن $\varepsilon_{\ell t}$ و $\eta_{\ell t}$ همانند مدل GP نرمال فرض می‌شوند و ρ پارامتر همبستگی زمانی بین $(-1, 1)$ است. اگر $\rho = 0$ باشد، مدل GP حاصل می‌شود. در این مدل مقادیر شروع $\mathbf{u}_\ell$ برای هر $\ell = 1, \dots, r$ از توزیع $N(\mu_\ell, \sigma_\ell^2 S_\ell)$ مشخص می‌شود، که در آن S_ℓ ماتریس همبستگی فضایی حاصل از تابع همبستگی ماترن با پارامترهای φ و ν همانند $\eta_{\ell t}$ است. بنابراین همه پارامترهای مدل به صورت $\theta = (\beta, \rho, \sigma_\varepsilon^2, \sigma_\eta^2, \varphi, \nu, \mu_\ell, \sigma_\ell^2)$ می‌باشند. با در نظر گرفتن پیشین معمول $\pi(\theta)$ برای پارامترها و با لگاریتم گرفتن توزیع پسین توأم پارامترها و مقادیر گمشده به صورت

$$\begin{aligned} \log \pi(\theta, \mathbf{u}_{\ell t}, \mathbf{z}_{\ell t}^* | \mathbf{z}_{\ell t}^{obs}) \propto & -\frac{N}{2} \log \sigma_\varepsilon^2 - \frac{1}{2\sigma_\varepsilon^2} \sum_{\ell=1}^r \sum_{t=1}^{T_\ell} (\mathbf{z}_{\ell t} - \mathbf{u}_{\ell t})' (\mathbf{z}_{\ell t} - \mathbf{u}_{\ell t}) \\ & - \frac{1}{2} \log |\sigma_\ell^2 S_\ell| - \frac{1}{2\sigma_\eta^2} \sum_{\ell=1}^r \sum_{t=1}^{T_\ell} (\mathbf{u}_{\ell t} - \mathbf{X}_{\ell t} \beta)' S_\eta^{-1} (\mathbf{u}_{\ell t} - \mathbf{X}_{\ell t} \beta) \\ & - \frac{\sum_{\ell=1}^r T_\ell}{2} \log |\sigma_\eta^2 S_\eta| \\ & - \frac{1}{2\sigma_\eta^2} \sum_{\ell=1}^r \sum_{t=1}^{T_\ell} (\mathbf{u}_{\ell t} - \mu_\ell)' S_\eta^{-1} (\mathbf{u}_{\ell t} - \mu_\ell) + \log \pi(\theta) \end{aligned}$$

به دست می‌آید (برای اطلاعات بیشتر به مراجع فینلی و همکاران (۲۰۰۹)، فینلی و همکاران (۲۰۱۲) مراجعه شود). در ادامه این دو مدل را روی داده‌های تلفات جاده ای استان قم برازش و مورد مقایسه قرار می‌دهیم.

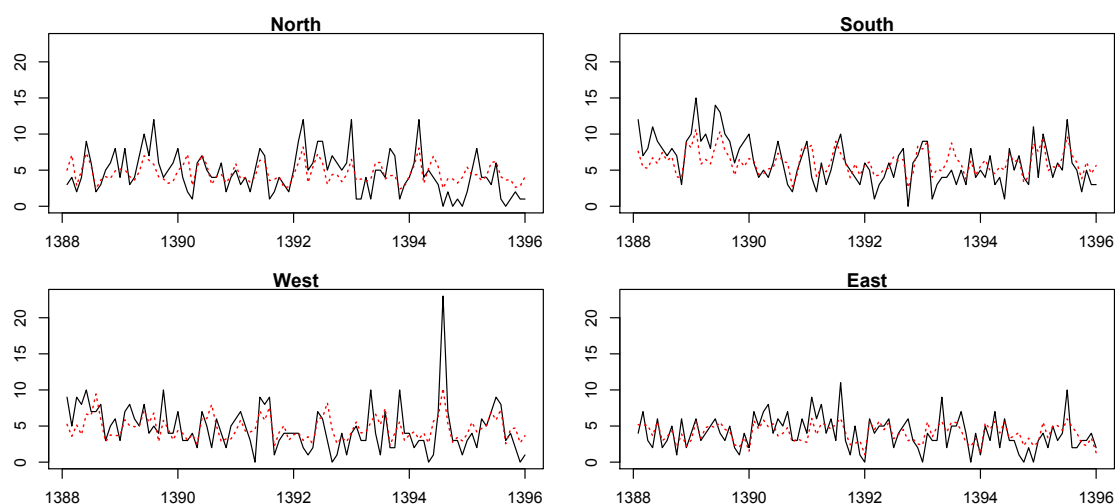
۱.۴ نتیجه تحلیل فضایی زمانی

لازم به ذکر است که موقعیت مکانی در مدل فضایی-زمانی برای چهار محور در دسترس بود، لذا از نظر مکانی تعداد کمی موقعیت جغرافیایی وجود دارد و با توجه به اینکه از نظر زمانی ۹۶ زمان است، بنابراین از نظر مکانی این مدل کارایی لازم را نسبت به زمان ندارد و این در مقدار $RMSE$ برای مدل فضایی-زمانی و مدل‌های سری زمانی چند متغیره ارائه شده در بخش ۳ کاملاً مشهود است. اگر بتوان در محورهای بیشتری با تفکیک ریزتری از نظر طول و عرض جغرافیایی مشاهدات را جمع‌آوری کرد، حتماً مدل‌های فضایی-زمانی ارائه شده در این بخش هم از نظر زمان اجرا محاسبات و هم قابل انجام بودن مدل مناسب‌تر از مدل‌های سری زمانی چند متغیره است.

جدول ۴: نتایج برازش مدل‌های فضایی زمانی GP و AR

	مدل GP			مدل AR		
	برآورد	انحراف معیار	فاصله اطمینان	برآورد	انحراف معیار	فاصله اطمینان
<i>Intercept</i>	-۰/۸۳۹۳	۰/۶۳۷۳	(-۲/۰۶, ۰/۴۴)	-۰/۸۲۵۱	۰/۶۴۱۴	(-۲/۰۷۵, ۰/۴۳۰)
<i>Accident</i>	۰/۰۲۰۹	۰/۰۱۰۴	(۰/۰۰۰۵, ۰/۰۴۰)	۰/۰۲۱۳	۰/۰۱۰۵	(۰/۰۰۰۸, ۰/۰۴۱)
<i>Injured</i>	۰/۰۲۲	۰/۰۲۰۶	(-۰/۰۱۸, ۰/۰۶۲)	۰/۰۲۱۹	۰/۰۲۰۱	(-۰/۰۱۸۱, ۰/۰۶۱)
<i>Riding</i>	۱/۲۶۵۹	۰/۰۴۳۶	(۱/۱۷۸, ۱/۳۵)	۱/۲۷	۰/۰۴۴۶	(۱/۱۸۶, ۱/۳۶۰)
σ_{ε}^2	۰/۰۱۰۱	۰/۰۰۰۸	(۰/۰۰۰۸۶, ۰/۰۱۱۸)	۰/۰۱۰۲	۰/۰۰۰۸	(۰/۰۰۰۸۶, ۰/۰۱۱۸)
σ_{η}^2	۵/۸۵۶۸	۰/۴۲۱۵	(۵/۰۹۴, ۶/۷۴)	۵/۸۸۲	۰/۴۳۷	(۵/۰۹۲, ۶/۸۰۴)
φ	۰/۰۱	۰/۰۰۰۱	(۰/۰۱, ۰/۰۲)	۰/۰۱۱	۰/۰۰۰۱	(۰/۰۱, ۰/۰۲۱)
ρ	-	-	-	-۰/۰۱۲۵	۰/۰۱۶۴	(-۰/۰۴۵, ۰/۰۱۸)
<i>RMSE</i>		۲/۱۳۲۱۹۸			۲/۱۲۱۸۴۲	

شکل ۸: نمودار سری زمانی که همان مرگ و میر به همراه برازش‌های مدل فضایی زمانی AR برای چهار محور استان قم را نمایش می‌دهد که خوبی برازش مدل AR به داده‌های واقعی را نشان می‌دهد.

شکل ۸: نمودار سری زمانی داده‌های مرگ و میر به همراه برازش مدل AR (خط‌چین) برای چهار محور استان قم

در اینجا متغیر پاسخ را تعداد مرگ و میر و متغیرهای کمکی را تعداد تصادفات، تعداد مجروحین و تعداد وسیله نقلیه سواری در نظر گرفته‌ایم. نتایج برازش مدل برای ۱۰۰۰۰ تکرار با مقادیر ذخیره یا سوخته ۵۰۰۰ در جدول ۴ خلاصه شده‌اند. با توجه به نتایج $RMSE$ جدول، مدل‌های فضایی زمانی ذکر شده تفاوت چندانی ندارند. با توجه به نتایج فاصله اطمینان ۹۵٪ برای جدول ۴، در بین متغیرهای کمکی هر دو مدل معنی‌داری دو متغیر کمکی تعداد تصادفات و وسیله نقلیه سواری را نشان می‌دهند (چون عدد صفر داخل فاصله اطمینان قرار ندارد).

۵ بحث و نتیجه‌گیری

با توجه به ماهیت داده‌های مورد مطالعه که به زمان و مکان وابسته هستند، در این مقاله به تحلیل زمانی و فضایی-زمانی داده‌ها پرداختیم. در یک تحلیل اکتشافی بر اساس مدل‌های خطی و خطی تعمیم یافته، متغیرهای کمکی تاثیر گذار بر نرخ

مرگ و میر تصادفات شناسایی شدند. سپس براساس نتایج حاصل از مدل‌های سری‌زمانی چندمتغیره، بهترین مدل براساس ملاک‌های MSE و AIC برای VARMA(۱،۱) است و در مدل‌های فضایی-زمانی بهترین مدل، مدل AR رتبه اول انتخاب گردید. در مقایسه کلی بین مدل‌های سری‌زمانی چندمتغیره و مدل‌های فضایی-زمانی، مدل‌های فضایی-زمانی عملکرد اجرایی محاسباتی مناسب‌تری برای این داده‌ها داشتند.

مراجع

- محمدزاده، م.، (۱۳۹۸)، *آمار فضایی و کاربردهای آن*، چاپ سوم، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران،
- والی، ح. (۱۳۹۲). تحقیق و بررسی علت وقوع حادثه درحوزه ۳۰ کیلومتری شهرها، انتشارات یاقوت.
- ویلیام دبلیو.اس.وی. ترجمه نیرومند ح.ع.، (۱۳۹۱) تحلیل سری‌های زمانی، روش‌های یک متغیری و چندمتغیری، انتشارات دانشگاه فردوسی مشهد.
- Cressie, N. (1993), *Statistics for Spatial Data*, John Wiley, New York.
- Cressie, NAC, and Wikle CK., (2011). *Statistics for Spatio-Temporal Data*. John Wiley & Sons, New York.
- Finley AO., Sang H., BANERJEE S.,and Gelfand A., (2009), Improving the performance of predictive process modeling for large datasets. *Computational Statistics and Data Analysis*, 53, 2873–2884.
- FINLEY, A., Banerjee, S., and Gelfand, A., (2012), Bayesian Dynamic Modeling for Large Space-Time Datasets Using Gaussian Predictive Processes. *Journal of Geographical Systems*, 14(1), 29-47.
- Gelfand, A., Banerjee, S., and Gamerman, D., (2005). Univariate and Multivariate Dynamic Spatial Modelling. *Environmetrics*, 16, 465-479.
- Gelfand, A.E., (2012), Hierarchical Modeling for Spatial Data Problems. *Spatial Statistics*, 1, 30-39.
- Matérn B., (1986), *Spatial Variation*. 2nd edition. Springer-Verlag, Berlin.



مروری بر طرح‌های نمونه‌گیری فضایی ساخته شده از دنباله هالتون

حسین ویسی‌پور^۱، محمد مرادی

گروه آمار، دانشگاه رازی

چکیده: در این مقاله به مروری بر کاربرد طرح‌های نمونه‌گیری پذیرش متعادل، نمونه‌گیری پذیرش متعادل اصلاح شده و نمونه‌گیری افرازشده تکراری هالتون می‌پردازیم. در این طرح‌ها، واحدهای نمونه با استفاده از دنباله هالتون استخراج می‌شوند. این طرح‌ها در جوامع پیوسته و گسسته چند بعدی مورد استفاده قرار می‌گیرند و واحدهای نمونه به وسیله این روش‌ها به هر دو صورت با احتمال برابر و احتمال نابرابر قابل استخراج می‌باشند.

واژه‌های کلیدی: احتمال شمول، برآوردگر هورویتز-تامپسون، به‌پخشی، دنباله هالتون، طرح نمونه‌گیری متعادل فضایی، همبستگی فضایی.

کد موضوع‌بندی ریاضی (۲۰۱۰): 62N01.

۱ مقدمه

در بررسی جوامع طبیعی و فضایی اطلاع از موقعیت واحدهای جامعه بسیار مهم است. در چنین حالتی واحدها بر اساس موقعیت جغرافیایی اندیس گذاری می‌شوند. در بررسی این جوامع واحدهای جامعه بر اساس مختصاتشان انتخاب می‌شوند. طرح‌های نمونه‌گیری که واحدهای نمونه را بوسیله انتخاب موقعیتشان استخراج می‌کنند را طرح‌های نمونه‌گیری فضایی می‌گویند. در جوامع فضایی معمولاً واحدهای نزدیک به هم تحت تاثیر همدیگر و عوامل مشترک همبسته می‌باشند. در نظر گرفتن این وابستگی از طریق انتخاب مکان‌های واحد جامعه به صورت نمونه می‌باشد (استیونس و اولسن ۲۰۰۴). استفاده از اطلاعات مکانی واحدهای نمونه باعث می‌شود که استنباط‌های دقیق‌تری در مورد کل فضای جامعه به دست آید (مک دونالد ۲۰۱۴). در صورت وجود همبستگی فضایی، هنگامی که یک واحد از نقاط همسایه انتخاب شود، انتخاب سایر نقاط همسایه کمک زیادی به شناسایی جامعه هدف نمی‌کند (هدایت، رائو و استوفکن ۱۹۸۸). بنابراین در چنین حالتی

^۱ نام و ایمیل ارائه دهنده مقاله: حسین ویسی‌پور، h.veisipour@gmail.com

طرح های نمونه گیری نیاز است که شانس انتخاب همزمان واحدهای نزدیک به هم کم و شانس انتخاب واحدهای همزمان واحدهای دورتر زیاد باشد. و طرح های در اولویت قرار می گیرند که واحدهای انتخاب شده در نمونه به خوبی در فضای جامعه پخش شوند. که استیونس و اولسن (۲۰۰۴) این طرح ها را طرح های فضایی متعادل نامیدند. معروفترین طرح های نمونه گیری فضایی متعادل عبارتند از روش موزایک بندی تصادفی (RTS)^۱ که توسط اورتون و استهمن (۱۹۹۳) معرفی گردید که یک طرح طبقه بندی تصادفی دو مرحله می باشد. در مرحله اول یک موزایک بندی با شبکه منظم به طور تصادفی بر روی دامنه واقع می شود و سپس در مرحله دوم یک نقطه تصادفی درون هر سلول موزایک انتخاب می گردد. استیونس و اولسن (۲۰۰۴) طرح موزایک بندی تصادفی تعمیم یافته (GRTS)^۲ را معرفی کردند که فضای دوبعدی را به یک بعدی تبدیل می کند و سپس از روش سیستماتیک یک بعدی با احتمال نابرابر استفاده کرده و واحدهای نمونه را انتخاب کردند. که انتخاب نمونه ها به صورت احتمال نابرابر و متناسب با طول نیز می تواند انجام گردد. باندسون و تربورن (۲۰۰۸) نمونه گیری پواسن همبسته (CPS)^۳ را برای انتخاب با احتمال نابرابر معرفی کردند. که به صورت مرحله ای اجرا شده و در هر مرحله احتمال شمول اصلاح می گردد. گفرستروم (۲۰۱۲) نمونه گیری پواسن همبسته فضایی (SCPS)^۴ را با ادغام کردن یک اندازه از فاصله ی واحدها در انتخاب کردن واحدها معرفی کرد که تعمیم روش (CPS) بود. گفرستروم و همکاران (۲۰۱۲) دو روش محلی محوری (LPM^۱, LPM^۲) را معرفی کردند، این روش ها تعمیم روش محوری (PM) می باشند که توسط دویل و تایلی (۲۰۱۲) معرفی شده است. این طرح ها برای انتخاب واحدهای نمونه π ps می باشند که نمونه در k مرحله انتخاب می گردد، در هر مرحله احتمال شمول برای هر دو واحد اصلاح می گردد و تا زمانی که احتمال شمول به یک یا صفر برسد این مراحل ادامه دارد. سه تا از جدیدترین روشهای نمونه گیری که از دنباله هالتون ساخته شده اند عبارتند از نمونه گیری فضایی متعادل پذیرفتنی (BAS)^۶ رابرتسون و همکاران (۲۰۱۳)، نمونه گیری فضایی متعادل پذیرفتنی اصلاح شده رابرتسون و همکاران (۲۰۱۷) و نمونه گیری فضایی افزایشده تکراری هالتون (HIP)^۷ که رابرتسون و همکاران (۲۰۱۸) آن را معرفی کردند.

در این مقاله به مروری بر مقالاتی می پردازیم که طرح های BAS، BAS اصلاح شده و HIP را معرفی کرده اند. در بخش دوم دنباله هالتون و خواص آن برگرفته از مقاله های هالتون (۱۹۶۰) و پرایس و پرایس (۲۰۱۲) را بیان می کنیم. در بخش سوم مقاله های ارائه شده توسط رابرتسون و همکاران (۲۰۱۳، ۲۰۱۷ و ۲۰۱۸) که به معرفی طرح های نمونه گیری BAS، BAS اصلاح شده و HIP و خواص آنها پرداخته اند را بررسی و نتایج را بیان می کنیم.

۲ دنباله هالتون

هالتون (۱۹۶۰) بردار d بعدی، $\{x_k\}_{k=0}^{\infty}$ در $[0, 1]^d$ را معرفی کرد که هر مولفه آن یک دنباله از اعداد شبه تصادفی است که بطور یکنواخت در بازه $[0, 1]$ توزیع می شود. مولفه i ام از هر نقطه دنباله به پایه b_i مرتبط است که $b_2 = 3, b_1 = 2$ و b_J برابر با J امین عدد اول طبیعی است. مقدار مولفه i از نقطه k ام این دنباله به صورت

$$x_k^{(i)} = \phi_{b_i}(k) = \sum_{j=0}^{\infty} \left\{ \left\lfloor \frac{k}{b_i^j} \right\rfloor \bmod b_i \right\} \frac{1}{b_i^{j+1}} \quad (1.2)$$

¹Random Tessellation Stratified

²Generalized Random Tessellation Stratified

³Correlated Poisson Sampling

⁴Spatially Correlated Poisson Sampling

⁵Local Pivotal Method

⁶Balanced Acceptance Sampling

⁷Halton Iterative Partitioning

محاسبه می‌شود، که در آن $[x]$ تابع جز صحیح است که برابر بزرگترین عدد صحیح کمتر یا مساوی x می‌باشد و $Mod k$ تابع باقیمانده است که باقیمانده تقسیم عدد k بر b را می‌دهد. به عنوان مثال، جمله سوم دنباله هالتون دوبعدی برابر $x_3 = (\phi_2(3), \phi_3(3)) = (3/4, 1/9)$ است.

۱.۲ خواص دنباله هالتون

خاصیت ۱. هالتون (۱۹۶۰) و پرایس و پرایس (۲۰۱۲) ثابت کردند برای هر عدد صحیح غیر منفی J_i و $B = \prod_{i=1}^d b_i^{J_i}$ از تعداد B نقطه متوالی از دنباله هالتون با پایه‌های b_i دقیقاً یک نقطه در هر خانه تعریف شده زیر قرار می‌گیرد.

$$\prod_{i=1}^d [m_i b_i^{-J_i}, (m_i + 1) b_i^{-J_i}) \quad (2.2)$$

که برای هر i ، $m_i = 0, 1, 2, \dots, b_i^{J_i}$ می‌باشد، که نشان می‌دهد دنباله هالتون به طور یکنواخت در بازه $[0, 1]^d$ پخش می‌گردد.

رابرتسون و همکاران (۲۰۱۸) این خانه‌ها را جعبه‌های هالتون نامیدند و این جعبه‌ها را با اندیس‌های $\{0, 1, \dots, B-1\}$ نامگذاری کردند. برای مطالعه بیشتر مقاله رابرتسون و همکاران (۲۰۱۸) را ببینید.

خاصیت ۲. پرایس و پرایس (۲۰۱۲) ثابت کردند، اگر X_k در یک خانه هالتون خاص قرار گیرد در آن صورت مقدار تابع باقیمانده k نسبت به B یکی از مقدارهای $\{0, 1, \dots, B-1\}$ است، بنابراین نقاط $X_k, X_{k+B}, X_{k+2B}, \dots$ از دنباله هالتون در همان جعبه خاص قرار می‌گیرند که نشان می‌دهد دنباله هالتون یک دنبای شبه دوه‌ای با دوره تناوب B می‌باشد.

۳ طرح‌های نمونه‌گیری فضایی

۱.۳ طرح نمونه‌گیری پذیرش متعادل

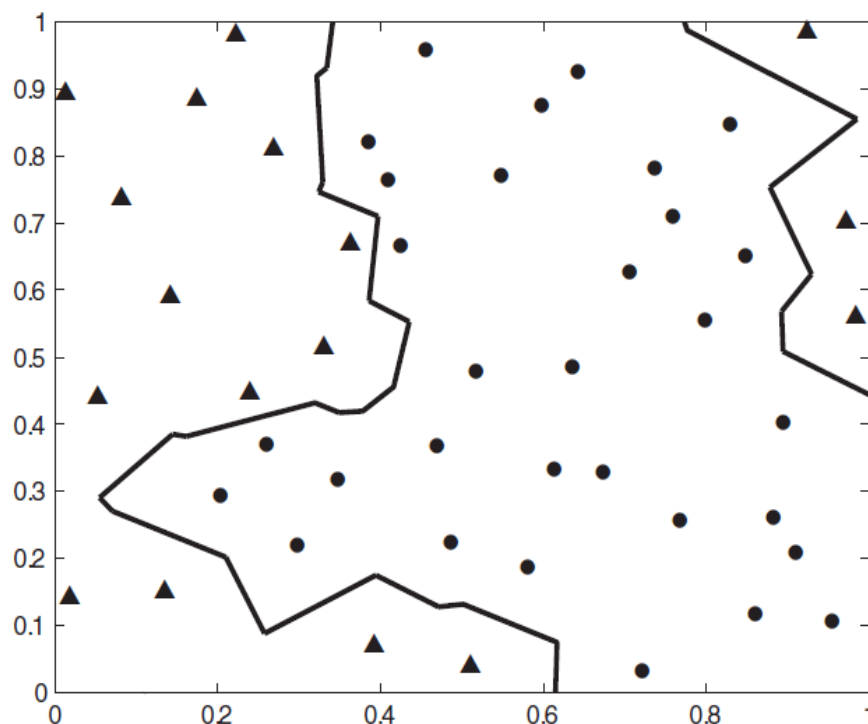
طرح نمونه‌گیری BAS توسط رابرتسون و همکاران (۲۰۱۳) معرفی گردید. که از دنباله شبه تصادفی هالتون با شروع تصادفی وانگ و هیکرنل (۲۰۰۰) برای به دست آوردن یک نمونه n تایی BAS استفاده کردند. یک نمونه ۳۰ تایی که به روش BAS انتخاب شده است در شکل ۱ نشان داده شده است. نقاط داخل ناحیه مورد بررسی با نماد $\bullet$ و نقاط خارج ناحیه با نماد $\blacktriangle$ نمایش داده شده‌اند. تعداد ۴۷ نقطه متوالی از دنباله هالتون با شروع تصادفی در انتخاب نمونه استفاده شده‌اند، که ۱۷ نقطه آن جهت انتخاب واحد در نمونه پذیرفته نشده‌اند. آنها بیان کردند که این طرح دارای خواص یک طرح خوب مانند سادگی اجرا، به‌پخشی، قابل استفاده در فضای پیوسته و گستره با ابعاد بالا، توانایی تغییر در اندازه نمونه در هر مرحله از اجرای طرح، قابل استفاده در جوامع با شکل نامنظم و نیز با احتمال نابرابر قابل استفاده می‌باشد.

با شبیه سازی واریانس برآوردگر هورویتز-تامپسون میانگین جامعه $\frac{y_i}{\pi_i}$ $\hat{\mu}_{HT} = \frac{1}{n} \sum_{i=1}^n$ طرح BAS با طرح‌های نمونه‌گیری تصادفی ساده (SRS) ^۸، GRTS، LPM۱، LPM۲، StrataSRS، ^۹ و SPCS مشاهده کردند که این واریانس در نمونه BAS کمتر از نمونه‌های به دست آمده از طرح‌های دیگر است. همچنین با شبیه سازی میزان به‌پخشی طرح BAS را با طرح‌های SRS و GRTS مشاهده کردند که به‌پخشی این طرح بهتراست.

^۸Simple Random Sampling

^۹Stratified Simple Random Sampling

یکی از عیب‌های این طرح این است که وقتی در انتخاب واحدهای نمونه، واحدهای بی پاسخی یا جعبه‌های بدون واحد وجود داشته باشد، احتمال شمول واحدهای به دست آمده از نمونه‌های BAS با احتمال شمول تعیین شده واحدها قبل از نمونه‌گیری یکسان نیست. این عیب اریبی و واریانس برآوردکننده را افزایش می‌دهد.



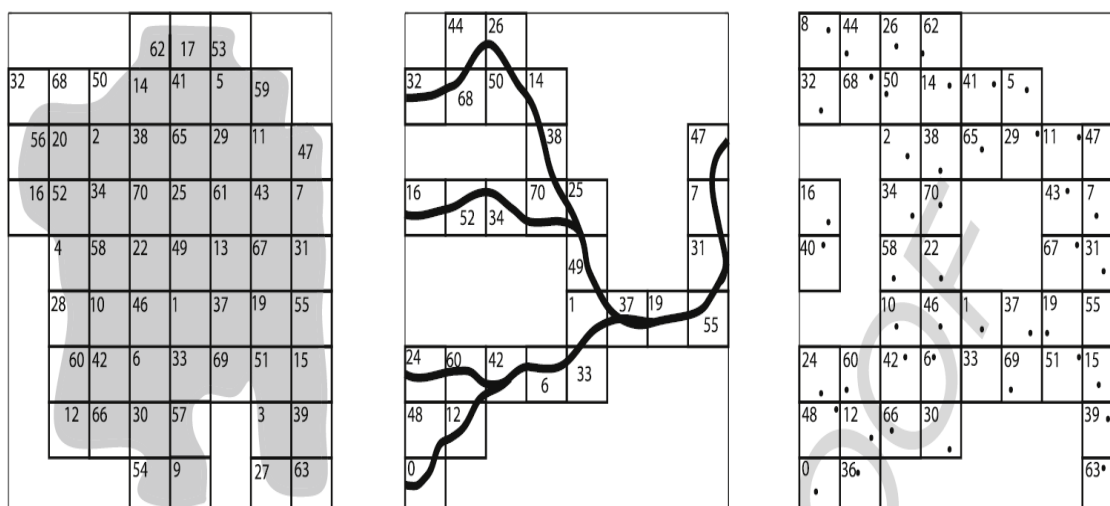
شکل ۱: نقاط داخل ناحیه مورد بررسی (●) و نقاط خارج ناحیه (▲).

۲.۳ طرح نمونه‌گیری پذیرش متعادل اصلاح شده

رابرتسون و همکاران (۲۰۱۷) عیب طرح نمونه‌گیری BAS را اصلاح کردند. آنها دو روش برای برطرف کردن این عیب پیشنهاد کردند، که روش اول برای نمونه‌گیری با احتمال برابر و نابرابر و روش دوم برای نمونه‌گیری با احتمال برابر استفاده می‌شود.

۱. اگر اولین نقطه دنباله هالتون با شروع تصادفی منجر به انتخاب واحد نمونه نشود در آن صورت این دنباله را نادیده می‌گیرند و سپس از یک دنباله دیگر هالتون با شروع تصادفی استفاده می‌کند.

۲. چون ممکن است تعداد زیادی دنباله منجر به انتخاب نشوند، لذا در صورت امکان از چارچوب هالتون استفاده می‌کند که در این حالت ابتدا تعداد $B = \prod_{i=1}^d b_i^{J_i}$ ، جعبه هالتون را در نظر می‌گیرند و سپس در جوامع گسسته با توجه به مکان واحدهای جامعه به هر واحد جامعه و در جوامع پیوسته با توجه به موقعیت ناحیه‌های افراز شده به هر ناحیه یک جعبه هالتون اختصاص می‌دهند. در ادامه به روش نمونه‌گیری تصادفی یک عدد از مجموعه $\{0, 1, \dots, B\}$ انتخاب و متناسب با عدد انتخاب شده، n جعبه هالتون متوالی به عنوان مکان‌های واحدهای نمونه مشخص می‌گردند (شکل ۲). شماره خانه‌های چارچوب هالتون برابر با باقیمانده تقسیم اندیس نقاط دنباله هالتون بر ۷۲ است.



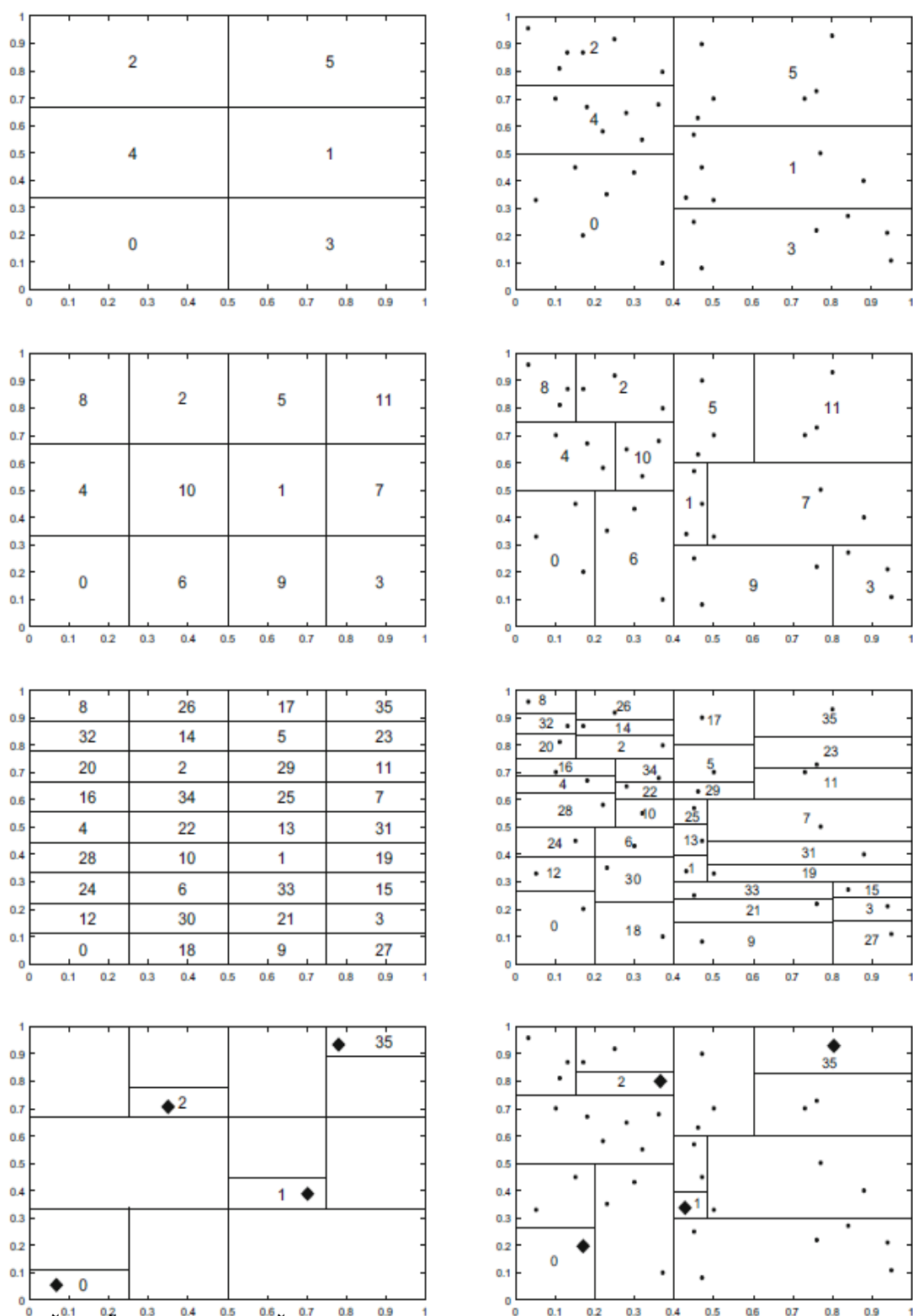
شکل ۲: چارچوب هالتون برای ناحیه مورد بررسی پیوسته (سایه دار)، خطی و نقطه‌ای با $b_1 = 2$ و $b_2 = 3$ ، $J = (3, 2)$.

۳.۳ طرح نمونه‌گیری افزایش شده تکراری هالتون

طرح نمونه‌گیری HIP توسط رابرتسون و همکاران (۲۰۱۳) معرفی گردید. که از دنباله شبه تصادفی هالتون با شروع تصادفی (وانگ و هیکرنل ۲۰۰۰) و یک افزایش تکراری هالتون از جامعه با تعداد ناحیه‌های برابر با B برای به دست آوردن یک نمونه n تایی HIP استفاده کردند. شکل ۳ سمت چپ، خانه‌های هالتون برای $B = 2 \times 3 = 6$ ، $B = 2 \times 3 = 12$ و $B = 2 \times 3^2 = 36$ را نشان می‌دهد. نقطه X_K از دنباله هالتون در خانه هالتون با شماره برابر با باقیمانده تقسیم K بر B قرار می‌گیرد. بعنوان مثال X_5 برای $B = 6$ در خانه شماره ۵، برای $B = 12$ در خانه شماره ۱۱ و برای $B = 36$ در خانه شماره ۲۳ قرار می‌گیرد. شکل سمت راست، افزایش‌های تو در تو برای یک جامعه نقطه‌ای با احتمال برابر و با $N = 36$ را نشان می‌دهد. به افزایش‌های انجام شده شماره متناظر با شماره خانه‌های هالتون اختصاص یافته است. ردیف آخر سمت چپ یک نمونه ۴ تایی HIP از یک ناحیه پیوسته با احتمال برابر را نشان می‌دهد که $K = 35$ و واحدهای نمونه در خانه‌های با شماره‌های $S_{35} = \{35, 0, 1, 2\}$ قرار می‌گیرند، و شکل سمت راست این نمونه ۴ تایی را در جامعه نقطه‌ای نشان می‌دهد.

این طرح دارای اکثر خاصیت‌های یک طرح فضایی متعادل خوب مانند سادگی اجرا، به پخشی، قابل استفاده در فضای پیوسته و گستره با ابعاد بالا، توانایی تغییر در اندازه نمونه در هر مرحله از اجرای طرح، قابل استفاده در جوامع با شکل نامنظم، قابل استفاده با احتمال نابرابر و نیز در صورت وجود واحدهای بی پاسخ احتمالی شامل به دست آمده از نمونه‌ها با احتمال شمول قبل از نمونه‌گیری خیلی نزدیک به هم می‌باشند.

در سه جامعه مصنوعی، شبیه‌سازی واریانس برآوردگر هورویتز-تامپسون نمونه‌های این طرح با طرح‌های نمونه‌گیری $GRTS$ ، SRS و LPM_2 انجام گرفته و مشاهده گردید که در اکثر مواقع این واریانس، در طرح HIP کمتر از طرح‌های دیگر است. و نیز با شبیه‌سازی میزان به پخشی نسبی این طرح با طرح‌های LPM_1 ، LPM_2 ، $GRTS$ و $GRTS(2n)$ مورد مقایسه قرار گرفته است. نتایج شبیه‌سازی نشان می‌دهد، این طرح و طرح LPM_2 به پخشی بهتری دارند.



شکل ۳: سمت چپ: خانه‌های هالتون برای $B = 2 \times 3 = 6$, $B = 2 \times 3 = 12$, و $B = 2^2 \times 3^2 = 36$ و سمت راست، افرازهای تو در تو برای یک جامعه نقطه‌ای با احتمال برابر و با $N = 36$.

مراجع

Bondesson, L., Thorburn, D. (2008), A List Sequential Sampling Method Suitable for Real-Time Sampling, *Canadian Journal of Statistic*, **35**, 466-483.

- Deville, J.C., Till'e, Y. (1998), Unequal Probability Sampling without Replacement Through a Splitting Method, *Biometrika*, **85**, 89-101.
- Grafström, A. (2012), Spatially Correlated Poisson Sampling, *Journal of Statistical Planning and Inference*, **142**, 139-147.
- Grafström, A., Lundström, NLP., Schelin, L. (2012), Spatially Balanced Sampling through the Pivotal Method, *Biometrics*, **68**, 514-520.
- Halton, J. H. (1993), On the Efficiency of Certain Quasi-Random Sequences of Points in Evaluating Multi-Dimensional Integrals, *Numerische, Mathematik*, **2**, 339-359.
- Hedayat, A. S., Rao, C. R., Stufken, J. (1988), Sampling Plans Excluding Contiguous Units, *Journal of Statistical Planning and Inference*, **19**, 159-170.
- McDonald T. (2014), Sampling Designs for Environmental Monitoring, In: *Manly BFJ, Navarro Alberto JA (eds) Introduction to Ecological Sampling. CRC Press, Taylor and Francis Group, Boca Raton*
- Overton, W. S., Stehman, S. V. (1993), Properties of Designs for Sampling Continuous Spatial Resources From a Triangular Grid, *Communications in Statistics, Part A Theory and Methods*, **22**, 2641-2660.
- Pric, C. J., Pric, C. J. (2012), Recycling Primes in Halton Sequences: An Optimization Perspective, *Advanced Modeling and Optimization*, **14**, 17-29.
- Robertson, B. L., Brown, J. A., McDonald, T., Jaksons, P. (2013), BAS: Balanced Acceptance Sampling of Natural Resources, *Biometrics*, **3**, 776-784.
- Robertson, BL., McDonald, T., Price, C. J., Brown, JA. (2017), A Modification of Malanced Acceptance Sampling. *Statistics & Probability Letters*, **129**, 107-112.
- Robertson, B. L., McDonald, T., Price, C. J., Brown, J. A. (2018), Halton Iterative Partitioning: Spatially Balanced Sampling via Partitioning, *Environmental and Ecological Statistics*, **25**, 305-323.
- Stevens, DL, Jr., Olsen, AR. (2004), Spatially Balanced Sampling of Natural Resources, *Journal of the American Statistical Association*, **99**, 262-278.
- Wang, X., Hickernell, F. J. (2000), Randomized Halton sequences, *Mathematical and Computer Modelling*, **32**, 887-899.



استفاده از رتبه‌ها در تحلیل ممیزی فضایی داده‌های چند متغیره

فائزه یزدانی^۱، مجید عظیم محسنی، مهناز خلفی
گروه آمار، دانشگاه گلستان

چکیده: در این مقاله به روش ناپارامتری تحلیل ممیزی فضایی داده‌های چندمتغیره پرداخته می‌شود. ابتدا برای داده‌های چندمتغیره از روش مقیاس‌بندی چندبعدی رتبه‌ای یا غیر متریک مختصات فضایی تخصیص داده می‌شود. سپس به روش تحلیل شباهت رتبه‌ای به نحوه تفکیک گروه‌ها از لحاظ داده‌های چندمتغیره پرداخته می‌شود. بر اساس یک مجموعه داده واقعی کارایی روش مورد ارزیابی قرار می‌گیرد.

واژه‌های کلیدی: داده‌های فضایی، مقیاس‌بندی چندبعدی، تحلیل شباهت.
کد موضوع‌بندی ریاضی (۲۰۱۰): 62H30, 62H11.

۱ مقدمه

تحلیل ممیزی به ارائه مرزبندی میان چند گروه از لحاظ مشاهدات چندمتغیره می‌پردازد. اما به دلیل بعد بالای مشاهدات، نمایش گرافیکی این مرزبندی‌ها عملاً امکان‌پذیر نمی‌باشد. روش مقیاس‌بندی چندبعدی یک روش کارا برای ارائه مختصات فضایی در ابعاد پایین برای داده‌های چندمتغیره فراهم می‌کند. علاوه بر آن این امکان را فراهم می‌سازد که بتوان از تحلیل‌های فضایی برای مشاهدات چندمتغیره استفاده کرد. روش مقیاس‌بندی چندبعدی معمولاً به دو روش متریک و غیر متریک صورت می‌گیرد. در روش متریک از ماتریس فاصله یا عدم شباهت داده‌ها در یافتن مختصات فضایی کمک گرفته می‌شود. وجود مشاهدات دور افتاده در مقادیر فاصله ممکن است انحراف زیادی را در نتایج ایجاد کند. به همین دلیل توابع یکنوا از فاصله که حساسیت کمتری به مشاهدات دور افتاده دارند مورد استفاده قرار می‌گیرد. مرسوم‌ترین تابع یکنوا رتبه‌ها می‌باشد که مبنای بسیاری از مطالعات ناپارامتری می‌باشد. استفاده از رتبه‌ها منجر به ارائه روش غیر متریک برای مقیاس‌بندی چندبعدی گردیده است. برای مطالعه بیشتر در خصوص مقیاس‌بندی چندبعدی می‌توان به [لژاندر و لژاندر \(۲۰۱۲\)](#)، [آگراوال و همکاران \(۲۰۰۷\)](#)، [کاکسون \(۱۹۸۲\)](#)، [کودی و اسپرینگر \(۲۰۱۸\)](#) و [گرفلمن و همکاران \(۲۰۱۸\)](#) مراجعه

^۱ نام و ایمیل ارائه دهنده مقاله: فائزه یزدانی، yazdanifaeze@yahoo.com

کرد. چه داده‌ها به طور فضایی ثبت شده باشند و چه با مقیاس بندی چندبعدی به آن‌ها مختصات فضایی تخصیص داده شود، می‌توان از آنالیز شباهت فضایی رتبه‌ای به نحوه تفکیک گروه‌ها پرداخت. به طور دقیق‌تر می‌توان برحسب فواصل مختلف به تفکیک گروه‌ها پرداخت و آزمون آماری تفکیک گروه‌ها را انجام داد. برای مطالعه بیشتر در خصوص آنالیز شباهت فضایی رتبه‌ای می‌توان به [لژاندر و لژاندر \(۲۰۱۲\)](#)، [کلارک \(۱۹۹۳\)](#) و [باتیگیگ و همکاران \(۲۰۱۴\)](#) مراجعه کرد. در بخش دوم به طور خلاصه به روش‌های به کار رفته در تحلیل ممیزی فضایی رتبه‌ای پرداخته می‌شود. در بخش سوم بر اساس یک سری داده واقعی به کارایی روش‌های رتبه‌ای در تفکیک گروه‌ها پرداخته می‌شود.

۲ روش‌ها

مقیاس بندی چندبعدی یک روش برای تخصیص مختصات فضایی در فضای با ابعاد کمتر به داده‌های چندمتغیره می‌باشد. الگوریتم این روش به دو صورت متریک و غیر متریک انجام می‌گیرد. در این روش ابتدا فاصله داده‌ها با یک معیار فاصله مانند فاصله اقلیدسی محاسبه می‌شود و ماتریس فاصله یا عدم شباهت $D = [d_{ij}]$ نقطه آغازین روش مقیاس بندی چندبعدی می‌شود.

فرض کنید $\mathbf{x}_1, \dots, \mathbf{x}_n$ مختصات فرضی داده‌های چندمتغیره در فضای p بعدی باشد. به طور معمول $p = 2$ یا $p = 3$ در نظر گرفته می‌شود. هدف مقیاس بندی چندبعدی متریک یافتن مختصات $\mathbf{x}_1, \dots, \mathbf{x}_n$ به طوری است که داشته باشیم:

$$d_{ij} \leq d_{kl} \iff \|\mathbf{x}_i - \mathbf{x}_j\|^2 \leq \|\mathbf{x}_k - \mathbf{x}_l\|^2$$

برای یافتن این مختصات معمولاً از مینیمم سازی توابع استرس استفاده می‌شود که نمونه‌ای از آن به صورت

$$Stress = \left(\frac{\sum_i \sum_j (d_{ij} - \|\mathbf{x}_i - \mathbf{x}_j\|^2)^2}{\sum_{i,j} d_{ij}^2} \right)^{\frac{1}{2}}$$

است. در مقیاس بندی چندبعدی غیر متریک به جای استفاده از d_{ij} از تابعی یکنوا از آن استفاده می‌شود که به نوعی این مقادیر را اصلاح کند. به عنوان مثال بتواند مشاهدات دور افتاده را کنترل کند. به طور معمول در این مقیاس بندی چندبعدی از رتبه فواصل به جای خود فواصل استفاده می‌شود. اگر داشته باشیم $r_{ij} = rank(d_{ij})$ در این صورت مختصات داده‌ها به گونه‌ای تعیین می‌شود که داشته باشیم

$$r_{ij} \leq r_{kl} \iff \|\mathbf{x}_i - \mathbf{x}_j\|^2 \leq \|\mathbf{x}_k - \mathbf{x}_l\|^2$$

این روش طی دو گام انجام می‌شود:

گام ۱: یک مقدار اولیه برای مختصات در نظر گرفته شود.

گام ۲: مقادیر گام اول آنقدر تغییر پیدا کند که تابع استرس به کمترین مقدار خود برسد.

تحلیل شباهت یک آزمون ناپارامتری برای تحلیل ممیزی فضایی محسوب می‌شود. آماره این آزمون بر اساس شباهت درون گروه‌ها و بین گروه‌ها ساخته می‌شود. به طور دقیق‌تر این آماره ثبات گروه بندی را در فواصل مختلف مورد آزمون قرار می‌دهد. این آماره بر اساس رتبه‌ها شکل گرفته است به همین دلیل آزمون فرض آن به صورت ناپارامتری می‌باشد. آماره آزمون تحلیل شباهت R در فواصل مختلف به صورت

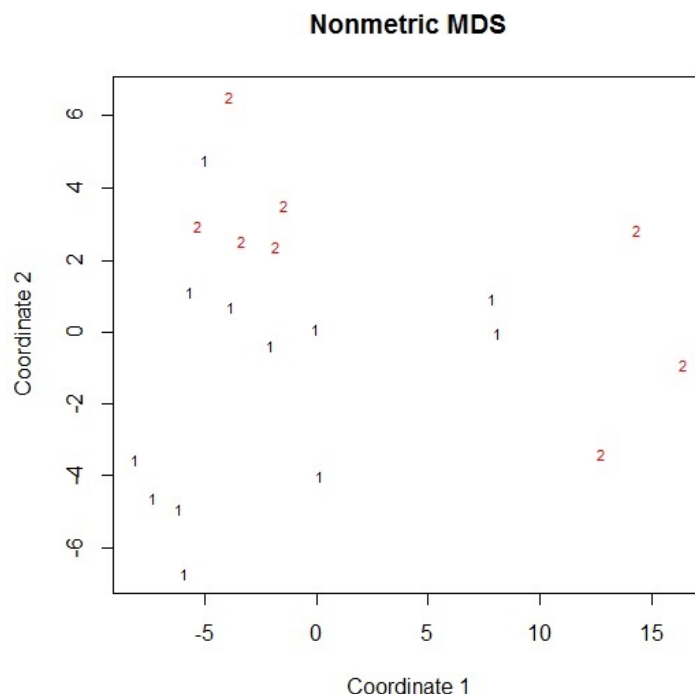
$$R(h) = \frac{r_B(h) - r_W(h)}{\frac{1}{4}n(n-1)} \quad (1.2)$$

ساخته می‌شود، که در آن $r_B(h)$ میانگین شباهت‌های رتبه‌ای بین گروه‌ها، $r_W(h)$ میانگین شباهت‌های رتبه‌ای درون گروه‌ها و n تعداد کل نمونه‌های مورد بررسی است. آماره R مقادیر خود را در بازه $(-1, 1)$ اختیار می‌کند که اعداد مثبت نشان‌دهنده شباهت بیشتر درون گروه‌ها است و مقادیر نزدیک به صفر نشان می‌دهد که هیچ تفاوتی بین گروه‌ها و درون گروه‌ها از نظر شباهت وجود ندارد. مقادیر منفی R نیز نشان‌دهنده شباهت بیشتر بین گروه‌ها نسبت به درون گروه‌ها است. محاسبه مقدار معنی‌داری به روش جایگزینی انجام می‌گیرد. الگوریتم محاسبه مقدار معنی‌داری از روش جایگزینی به‌صورت زیر است:

- گام ۱: مشاهدات در گروه‌ها به‌طور تصادفی جابه‌جا شود.
- گام ۲: آماره R برای مشاهدات گام اول محاسبه شود.
- گام ۳: گام اول و گام دوم B بار تکرار می‌شود مقدار آماره محاسبه و $R_1, \dots, R_B$ نام‌گذاری شوند.
- گام ۴: در صد R_i که از R کوچکتر هستند محاسبه می‌شود و به عنوان مقدار معنی‌داری در نظر گرفته شود.

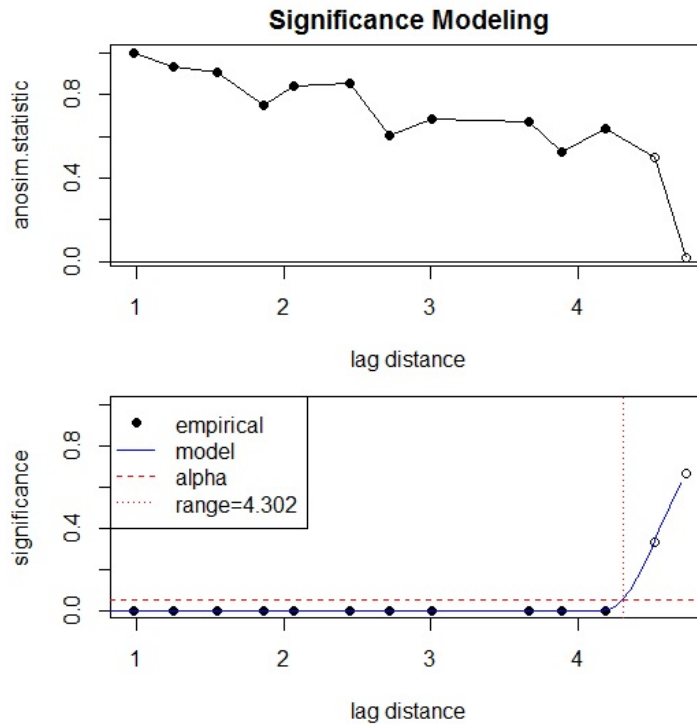
۳ نتایج عددی

در این بخش روش‌های ذکر شده در بخش دوم را برای یک مجموعه داده واقعی با دو گروه به کار می‌بریم. از دو گروه با وزن کم (کد ۱) و وزن زیاد (کد ۲) ۱۱ فاکتور خونی، PH ، تعداد رنگدانه کراتنن، ضریب اصلاح شده کراتنن، فسفات، کلسیم، فسفر، کراتنن، بور، کلرید، کولین، مس ثبت شده است. در این آزمایش ۱۲ نفر با وزن کم و ۸ نفر با وزن زیاد انتخاب شده‌اند. شکل ۱ آرایش فضایی مشاهدات چندمتغیره را در فضای دو بعدی نشان می‌دهد. تجمع مشاهدات با کد



شکل ۱: مختصات فضایی مشاهدات فاکتورهای خونی دو گروه با وزن کم (کد ۱) و وزن زیاد (کد ۲) با روش مقیاس بندی چند بعدی رتبه‌ای

یکسان در کنار هم نشان‌دهنده کارایی روش مقیاس‌بندی رتبه‌ای چندبعدی می‌باشد.



شکل ۲: نتایج حاصل از تحلیل شباهت فضایی: نمودار بالا مقدار آماره R در فواصل مختلف. نمودار پایین مقادیر معنی داری برحسب فاصله به همراه نمودار گوسی برازش داده شده

شکل ۲ نتایج حاصل از آنالیز شباهت داده‌ها برای دو گروه افراد با وزن زیاد و کم را نشان می‌دهد. بر اساس نمودار قسمت بالا مشخص می‌شود که در فواصل کمتر میزان شباهت درون گروه‌ها بالا و اختلاف بین گروه‌ها زیاد می‌باشد. هرچه فاصله بیشتر شده است عدم شباهت بین گروه‌ها کاهش پیدا کرده است. نمودار پایین مقادیر معنی داری در فواصل مختلف را براساس آزمون جایگشتی نشان می‌دهد. براساس این نمودار می‌توان قضاوت کرد که تا فاصله $h = 4/302$ اختلاف معنی داری بین گروه‌ها وجود دارد و پس از آن این فاصله دیگر اختلاف بین گروه‌ها معنی دار نیست. با توجه به این که این نمودار به صورت غیرنزولی افزایش می‌یابد می‌توان به آن یک تابع گوسی تجمعی به صورت

$$y^s(h) = C_h(1 - e^{C_w(h-C_c)^2}) \quad (1.3)$$

برازش داد، که در آن C_h ارتفاع منحنی، C_w عرض یا شیب منحنی و C_c مرکز منحنی است. در شکل ۲ قسمت پایین منحنی گاوسی برازش داده شده به مقادیر معنی داری نیز ترسیم شده است. پارامترهای برآورد شده در این مدل $Ch = 1/00534$, $Cc = 4/183173$, $Cw = 3/654079$ هستند.

بحث و نتیجه‌گیری

در این مقاله به یک تحلیل ممیزی فضایی برای داده‌های چندمتغیره پرداخته شد. روش مقیاس‌بندی چندبعدی این امکان را فراهم ساخت که نه تنها بتوان به داده‌های کلاسیک چندگانه آرایش فضایی داد، بلکه بتوان از آنالیز شباهت فضایی برای این

داده‌ها استفاده کرد. کلیه مراحل این تحلیل ممیزی فضایی مبتنی بر رتبه‌ها ارائه گردید. نتایج عددی که بر اساس مجموعه‌ای از داده‌های واقعی ارائه گردید نشان داد که استفاده از رتبه‌ها به‌خوبی، توانایی تفکیک گروه‌ها را دارد. نتایج عددی این مقاله براساس تحلیل ممیزی دو گروه از افراد ارائه گردید، اما این روش را به‌راحتی می‌توان برای چند گروه از افراد نیز به‌کار گرفت.

مراجع

Agarwal, S., Wills, J., Cayton, L., Lanckriet, G., Kriegman, D., and Belongie, S. (2007), Generalized Non-Metric Multidimensional Scaling, *In Artificial Intelligence and Statistics*, 11-18.

Buttigieg, Pier Luigi; Ramette, Alban (2014), A Guide to Statistical Analysis in Microbial Ecology: A Community-Focused, Living Review of Multivariate Data Analyses, *FEMS Microbiology Ecology*, 90, 543–550.

Clarke, K. R. (1993), Non-Parametric Multivariate Analyses of Changes in Community Structure, *Australian Journal of Ecology*, 18, 117-143.

Coxon, A. P. M. (1982), The User's Guide to Multidimensional Scaling, *London: Heinemann Educational Books*.

Ding, C. S. (2018), Fundamentals of Applied Multidimensional Scaling For Educational and Psychological Research, *Springer: Cham*.

Graffelman, J., Femenia, I. G., de Cid, R., and Barcelo-i-Vidal, C. (2018), Multidimensional Scaling and Relatedness Research, *bioRxiv*, 297879.

Legendre, P., and Legendre, L. F. (2012), *Numerical Ecology*, Elsevier, Vol. 24.