

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ



دانشکده علوم ریاضی

رشته آمار، گرایش آمار ریاضی

پایان نامه کارشناسی ارشد

یک مدل منعطف برای پیش‌گویی فضایی پاسخ‌های نرخ به کمک توزیع بتا-دوجمله‌ای

نگارنده: لیلا عابدین پور لیارجدمه

استاد راهنما

دکتر حسین باغیشنی

استاد مشاور

دکتر نگار اقبال

بهمن ۱۳۹۶



مدیریت تحصیلات تکمیلی

باسمه تعالی

شماره:
تاریخ:

فرم شماره (۳) صورتجلسه نهایی دفاع از پایان نامه دوره کارشناسی ارشد

با نام و یاد خداوند متعال، ارزیابی جلسه دفاع از پایان نامه کارشناسی ارشد خانم لیلا عابدین پور لیارجدمه با شماره دانشجویی ۹۴۱۱۷۲۴ رشته آمار گرایش آمار ریاضی تحت عنوان یک مدل منعطف برای پیش‌گویی فضایی پاسخ‌های نرخ به کمک توزیع بتا-دوجمله‌ای که در تاریخ ۹۶/۱۱/۹ با حضور هیأت محترم داوران در دانشگاه صنعتی شاهرود برگزار گردید به شرح ذیل اعلام می‌گردد:

قبول (با درجه: عالی) ☒ مردود ☐ نوع تحقیق: نظری ☒ عملی ☐			
عضو هیأت داوران	نام و نام خانوادگی	مرتبه علمی	امضاء
۱- استاد راهنمای اول	دکتر حسین باغیشنی	استادیار	
۲- استاد راهنمای دوم			
۳- استاد مشاور	دکتر نگار اقبال	استادیار	غایب
۴- نماینده تحصیلات تکمیلی	دکتر احمد نزاکتی رضازاده	دانشیار	
۵- استاد ممتحن اول	دکتر داود شاهسونی	دانشیار	
۶- استاد ممتحن دوم	دکتر محمد آرشی	دانشیار	

نام و نام خانوادگی رئیس دانشکده: دکتر ابراهیم هاشمی

تاریخ و امضاء و مهر دانشکده: ۹۶/۱۱/۹
تبصره: در صورتی که کسی مردود شود حداکثر یکبار دیگر (در مدت مجاز تحصیل) می‌تواند مجدداً دفاع نماید (دفاع مجدد نباید زودتر از ۴ ماه برگزار شود).



تقدیم به:
پدر و مادر عزیزم

به پاس عاطفه سرشار و گرمای امید بخش وجودشان
و خواهر و برادر مهربانم

به پاس محبت های بی دریغ شان
و همه آنان که اندیشیدن راه من آموختند، نه اندیشه مرا

سپاس‌گزاری...

سپاس خدای را که سخنوران، در ستودن او بمانند و شمارندگان، شمردن نعمت‌های او ندانند و کوشندگان، حق او را گزاردن نتوانند. سلام و دورد بر محمد و خاندان پاک او، طاهران معصوم، هم آنان که وجودمان وام‌دار وجودشان است و نفرین پیوسته بر دشمنان ایشان تا روز رستاخیز...

بدون شک جایگاه و منزلت معلم، اجل از آن است که در مقام قدردانی از زحمات بی‌شائبه او، با زبان قاصر و دست ناتوان، چیزی بنگاریم. اما از آن جایی که تجلیل از معلم، سپاس از انسانی است که هدف و غایت آفرینش را تامین می‌کند، وظیفه خود می‌دانم از زحمات بی‌دریغ استاد راهنمای خود، جناب آقای دکتر حسین باغیشنی که همواره با شکیبایی و درایت فراوان، مرا مستفیض راهنمایی‌های بی‌دریغ خود از لحاظ علمی و اخلاقی نمودند، صمیمانه تشکر و قدردانی نمایم که قطعاً بدون راهنمایی‌های ارزنده ایشان این مجموعه به انجام نمی‌رسید. از استاد گرامی، سرکار خانم دکتر نگار اقبال، که زحمت مشاوره این پایان‌نامه را تقبل فرمودند، کمال تشکر را دارم.

از اساتید محترم آقایان دکتر شاهسونی، دکتر آرشی، دکتر ربیعی و دکتر نزاکی، که در زمان تحصیل از محضرشان سود برده و از تجربیات ایشان استفاده کردم، سپاسگزاری می‌نمایم و امیدوارم همواره در پناه لطف الهی محفوظ و موفق باشند. همچنین تشکر و قدردانی فراوان خود را از خانواده عزیزم که با حمایت‌های خویش، همواره مشوق من بوده‌اند ابراز می‌نمایم.

در پایان نیز از دانشگاه علوم پزشکی گیلان و مرکز آموزشی درمانی رازی رشت جهت ارائه داده‌های مربوط به سرطان معده تشکر و قدردانی می‌کنم.

لیلا عابدین پور لیارجدمه

بهمن ۱۳۹۶

تعهد نامه

اینجانب لیلا عابدین پور لیارجدمه دانشجوی کارشناسی ارشد رشته آمار علوم ریاضی دانشگاه صنعتی شاهرود، نویسنده پایان نامه با عنوان یک مدل منعطف برای پیش‌گویی فضایی پاسخ‌های نرخ به کمک توزیع بتا- دوجمله‌ای، تحت راهنمایی دکتر حسین باغیثنی متعهد می‌شوم:

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش‌های دیگر پژوهش‌گران، به مرجع مورد استفاده استناد شده است.
- مطالب این پایان نامه، تا کنون توسط خود، یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ‌جا ارایه نشده‌اند.
- حقوق معنوی این اثر، به دانشگاه صنعتی شاهرود تعلق دارد، و مقالات مستخرج با نام “دانشگاه صنعتی شاهرود” یا “Shahrood University of Technology” به چاپ خواهند رسید.
- حقوق معنوی تمام افرادی که در به‌دست آوردن نتایج اصلی پایان نامه تاثیرگذار بوده‌اند، در مقالات مستخرج از پایان نامه رعایت می‌گردد.
- در تمام مراحل انجام این پایان نامه، در مواردی که از موجود زنده (یا بافت‌های آن‌ها) استفاده شده است، ضوابط و اصول اخلاقی رعایت شده‌اند.
- در تمام مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته (یا استفاده شده است)، اصل رازداری و اصول اخلاق انسانی رعایت شده است.

لیلا عابدین پور لیارجدمه

بهمن ۱۳۹۶

مالکیت نتایج و حق نشر

- تمام حقوق معنوی این اثر و محصولات آن (مقالات مستخرج، کتاب، برنامه‌های رایانه‌ای، نرم‌افزارها و تجهیزات ساخته‌شده) متعلق به دانشگاه صنعتی شاهرود می‌باشد. این مطلب باید به نحو مقتضی، در تولیدات علمی مربوطه ذکر شود.
- استفاده از اطلاعات و نتایج موجود در این پایان نامه بدون ذکر منبع مجاز نیست.

چکیده

مدلی که به‌طور معمول برای برازش پاسخ‌های شمارشی فضایی در جوامع متناهی مورد استفاده قرار می‌گیرد، مبتنی بر توزیع دوجمله‌ای است. در مباحث کاربردی، به‌دلیل وجود بیش‌پراکنش معمول در داده‌ها، استفاده از این پذیره توزیعی چندان مناسب نیست و می‌تواند به استنباط ناکارا در برآورد پارامترها و پیش‌گویی فضایی منتهی شود. به‌منظور رفع نقایص مدل دوجمله‌ای، یک رهیافت جانشین را برای مدل‌سازی نسبت‌های همبسته فضایی مبتنی بر توزیع بتا-دوجمله‌ای معرفی می‌کنیم. انعطاف‌پذیری توزیع بتا-دوجمله‌ای می‌تواند آن را به گزینه منتخب برای برازش پاسخ‌های بیش‌پراکنده در این جوامع تبدیل کند. در این پایان‌نامه، به پیش‌گویی فضایی با استفاده از مدل منعطف بتا-دوجمله‌ای و تعمیم آن برای تحلیل فضایی داده‌های شبکه‌ای در چارچوب استنباط بیزی می‌پردازیم. همچنین با مطالعه شبیه‌سازی، کارایی آن را نمایش می‌دهیم. برای نمایش عملکرد این مدل‌ها، دو مجموعه داده بارندگی استان سمنان و تعداد افراد مبتلا به سرطان معده در استان گیلان را تحلیل می‌کنیم.

کلمات کلیدی: داده‌های شمارشی فضایی، فرآیند بتا-دوجمله‌ای، بیش‌پراکنش، پیش‌گویی فضایی.

فهرست مقاله‌های مستخرج از پایان‌نامه

۱. عابدین پور، ل.، باغیشنی، ح و اقبال، ن.، (۱۳۹۶)، ”پیش‌گویی فضایی در فرآیندهای بتا-دوجمله‌ای“، مجموعه مقالات یازدهمین سمینار احتمال و فرآیندهای تصادفی، ۸۷-۹۷، قزوین.
۲. عابدین پور، ل.، باغیشنی، ح و اقبال، ن.، (۱۳۹۶)، ”پیش‌گویی فضایی پاسخ‌های نرخ: مقایسه مدل‌های کریگینگ بتا-دوجمله‌ای و هم‌کریگینگ دوجمله‌ای“، مجموعه مقالات دومین سمینار آمار فضایی و کاربردهای آن، (۱۶۷-۱۷۹)، شاهرود.
۳. عابدین پور، ل.، باغیشنی، ح و اقبال، ن.، (۱۳۹۶)، ”پیش‌گویی نرخ بارندگی در استان سمنان با مدل بتا-دوجمله‌ای فضایی“، مجموعه مقالات دومین کنفرانس ملی و اولین کنفرانس بین‌المللی محاسبات نرم، (۱۰۸۷-۱۰۸۱)، رودسر.
۴. مدل اتوبتا-دوجمله‌ای برای تحلیل نرخ سرطان معده در استان گیلان، مجله دانشگاه علوم پزشکی گیلان، ارسال شده.

فهرست مطالب

ق فهرست تصاویر

ش فهرست جداول

۱	۱ داده‌های شمارشی همبسته فضایی
۱	۱.۱ مقدمه
۴	۱.۱.۱ داده‌های فضایی
۵	۲.۱.۱ مدل‌های آماری فضایی: میدان تصادفی
۷	۳.۱.۱ ساختار همبستگی فضایی
۱۰	۲.۱ پیش‌گویی فضایی
۱۱	۱.۲.۱ کریگینگ معمولی
۱۴	۲.۲.۱ کریگینگ ساده
۱۴	۳.۲.۱ کریگینگ عام
۱۴	۳.۱ مدل‌های خطی تعمیم‌یافته
۱۶	۱.۳.۱ مدل‌های آمیخته خطی تعمیم‌یافته
۱۷	۲.۳.۱ مدل‌های آمیخته خطی تعمیم‌یافته فضایی
۱۷	۳.۳.۱ رگرسیون دوجمله‌ای
۱۸	۴.۳.۱ رگرسیون پواسون
۱۹	۴.۱ کریگینگ پواسون
۲۰	۱.۴.۱ مشخصات مدل
۲۲	۲.۴.۱ تغییرنگار کریگینگ پواسون
۲۲	۳.۴.۱ معادلات کریگینگ پواسون
۲۳	۵.۱ توزیع‌ها
۲۴	۱.۵.۱ توزیع دوجمله‌ای
۲۴	۲.۵.۱ توزیع بتا

۲۵	۲	مدل بتا-دوجمله‌ای فضایی
۲۵	۱.۲	مقدمه
۲۶	۲.۲	توزیع بتا-دوجمله‌ای
۲۸	۱.۲.۲	میانگین و واریانس توزیع بتا-دوجمله‌ای
۲۹	۲.۲.۲	توجیه جعبه و مهره برای مدل بتا-دوجمله‌ای
۳۰	۳.۲.۲	فرآیند فضایی بتا-دوجمله‌ای
۳۳	۳.۲	برآوردگر تغییرنگار بتا-دوجمله‌ای
۳۴	۴.۲	معادلات کریگینگ بتا-دوجمله‌ای
۴۱	۳	دقت برآوردیابی و پیش‌گویی مدل بتا-دوجمله‌ای فضایی
۴۱	۱.۳	مقدمه
۴۱	۱.۱.۳	روش تبدیل معکوس برای شبیه‌سازی داده‌ها
۴۳	۲.۱.۳	شبیه‌سازی میدان فضایی بتا-دوجمله‌ای
۴۴	۳.۱.۳	مشخصات داده‌های شبیه‌سازی شده بتا-دوجمله‌ای فضایی
۴۷	۴.۱.۳	برآورد پارامترهای توزیع بتا
۵۰	۵.۱.۳	ویژگی‌های تغییرنگار بتا-دوجمله‌ای
۵۲	۲.۳	روش‌های مختلف برآورد تغییرنگار
۵۸	۳.۳	برآورد پارامترهای ساختار وابستگی فضایی
۶۲	۴.۳	دقت پیش‌گویی
۶۷	۴	کارایی پیش‌گویی مدل بتا-دوجمله‌ای فضایی
۶۷	۱.۴	هم‌کریگینگ دوجمله‌ای
۶۹	۱.۱.۴	معادلات هم‌کریگینگ دوجمله‌ای
۷۰	۲.۱.۴	برآورد پارامترهای هم‌کریگینگ دوجمله‌ای
۷۲	۳.۱.۴	پیش‌گویی فضایی
۷۴	۲.۴	پیش‌گویی و نقشه پهنه‌بندی نرخ بارش در استان سمنان
۷۹	۵	تحلیل فضایی داده‌های شمارشی شبکه‌ای با مدل اتوبتا-دوجمله‌ای
۸۰	۱.۵	مقدمه
۸۱	۲.۵	داده‌های فضایی شبکه‌ای
۸۳	۱.۲.۵	روش‌های انتساب وزن‌های همسایگی
۸۴	۳.۵	مدل‌بندی داده‌های شبکه‌ای
۸۵	۴.۵	معرفی مدل‌های اتورگرسیو شرطی
۸۵	۱.۴.۵	زنجیرهای مارکوف
۸۶	۲.۴.۵	معرفی مدل

۳.۴.۵	بررسی بیش تر مدل های شرطی	۸۷
۵.۵	توزیع های نمایی شرطی	۸۹
۶.۵	توزیع بتا- دوجمله ای بازپارامتری شده	۹۳
۱.۶.۵	مدل اتوبتا- دوجمله ای	۹۴
۷.۵	استنباط بیزی تقریبی	۹۵
۱.۷.۵	مدل های گاوسی پنهان	۹۵
۲.۷.۵	تقریب لاپلاس آشیانه ای جمع بسته	۹۶
۸.۵	تحلیل داده های سرطان معده در استان گیلان	۹۷
۱.۸.۵	نرخ سرطان معده در استان گیلان	۹۹
۹.۵	نتیجه گیری و آینده تحقیق	۱۰۲
۱۰۵	مراجع	
۱۱۱	آ دستورات نرم افزار R	

فهرست تصاویر

۴	محل سکونت افراد مبتلا به سرطان معده در استان گیلان	۱.۱
۵	محل تخلقات رانندگی در یک مسابقه رالی	۲.۱
۸	یکی از حالت‌های تابع تغییرنگار	۳.۱
۴۲	میدان‌های شبیه‌سازی‌شده با روش تبدیل معکوس	۱.۳
۴۴	توابع چگالی بتا	۲.۳
۴۷	میانگین و واریانس شبیه‌سازی‌شده	۳.۳
۵۳	اثر قطعه‌ای نظری توزیع‌های بتا	۴.۳
۵۵	برآورد میانه و انحراف معیار پارامترهای تغییرنگار برای ۵۴ سناریو شبیه‌سازی بر اساس ۱۰۰ مجموعه داده شبیه‌سازی با استفاده از روش اول	۵.۳
۵۶	برآورد میانه و انحراف معیار پارامترهای تغییرنگار برای ۵۴ سناریو شبیه‌سازی بر اساس ۱۰۰ مجموعه داده شبیه‌سازی با استفاده از روش دوم	۶.۳
۵۷	برآورد میانه و انحراف معیار پارامترهای تغییرنگار برای ۵۴ سناریو شبیه‌سازی بر اساس ۱۰۰ مجموعه داده شبیه‌سازی با استفاده از روش سوم	۷.۳
۶۰	اثر قطعه‌ای فضایی برآورد شده	۸.۳
۶۰	ازاره فضایی برآورد شده	۹.۳
۶۱	دامنه فضایی برآورد شده	۱۰.۳
۶۴	میانگین خطای پیش‌گویی	۱۱.۳
۶۴	میانگین توان دوم خطای پیش‌گویی	۱۲.۳
۶۵	میانگین واریانس پیش‌گویی	۱۳.۳
۶۵	میانگین مقادیر پیش‌گویی	۱۴.۳
۶۶	همبستگی بین توزیع بتا پیش‌گویی و توزیع بتا پنهان	۱۵.۳
۶۶	واریانس مقادیر پیش‌گویی	۱۶.۳
۷۱	برآورد پارامترهای تغییرنگار بتا-دوجمله‌ای	۱.۴
۷۲	برآورد پارامترهای تغییرنگار دوجمله‌ای	۲.۴

۳.۴	مقایسه مقادیر MSPE برای دو روش کریگینگ بتا- دوجمله‌ای و هم کریگینگ
۷۳	دوجمله‌ای
۴.۴	میانگین توان دوم خطای پیش‌گویی دو مدل به ازای ۵۴ سناریوی شبیه‌سازی
۵.۴	موقعیت ایستگاه‌های هواشناسی استان سمنان
۶.۴	موقعیت ایستگاه‌های هواشناسی استان سمنان
۷.۴	نیم‌تغییرنگار تجربی و برازش مدل گاوسی به داده‌ها
۸.۴	نقشه پهنه‌بندی پیش‌گویی نرخ بارش در استان سمنان
۹.۴	واریانس نرخ پیش‌گویی بارش در استان سمنان
۱.۵	نمودار پراکنش مقادیر برازش‌شده در مقابل مقادیر واقعی نرخ سرطان معده
۱۰۱	در دو مدل دوجمله‌ای (سمت راست) و بتا- دوجمله‌ای (سمت چپ) . .
۲.۵	نقشه‌های پهنه‌بندی نرخ سرطان معده برای مدل‌های دوجمله‌ای (سمت
۱۰۲	راست) و بتا- دوجمله‌ای (سمت چپ)

فهرست جداول

۴۵	 میانگین و واریانس توزیع‌های بتا	۱.۳
۴۶	 سناریوهای مختلف مطالعه شبیه‌سازی	۲.۳
۴۸	 برآورد پارامترهای شکل $\hat{\alpha}$ و $\hat{\beta}$ برای هر یک از توزیع‌های بتا با دامنه ۵	۳.۳
۴۹	 برآورد پارامترهای شکل $\hat{\alpha}$ و $\hat{\beta}$ برای هر یک از توزیع‌های بتا با دامنه ۱۰	۴.۳
۴۹	 برآورد پارامترهای شکل $\hat{\alpha}$ و $\hat{\beta}$ برای هر یک از توزیع‌های بتا با دامنه ۱۵	۵.۳
	۶.۳ واریانس واقعی نسبت‌های نمونه‌ای دوجمله‌ای تحت توزیع‌های بتا در شرایط	
۵۱	 شبیه‌سازی با $n=3$	
۵۲	 پارامترهای فضایی واقعی برای نسبت نمونه‌ای $\frac{Z}{n}$ با اندازه نمونه برابر	۷.۳
۸۲	 شبکه منظم با پوشش شبکه‌ای از مربع‌ها	۱.۵
	۲.۵ نتایج برازش دو مدل بتا-دوجمله‌ای و دوجمله‌ای برای داده‌های سرطان	
۱۰۱	 معده	

فصل ۱

داده‌های شمارشی همبسته فضایی

۱.۱ مقدمه

در بسیاری از علوم مانند علوم محیطی، پزشکی، ژنتیک و زمین‌شناسی با داده‌هایی مواجه می‌شویم که مستقل از یکدیگر نیستند و وابستگی آن‌ها ناشی از موقعیتشان در فضای مورد مطالعه است، به‌طوری که ساختار وابستگی آن‌ها تابعی از فاصله مشاهدات از یکدیگر است. این‌گونه داده‌ها، داده‌های فضایی^۱ نامیده می‌شوند. به‌عنوان مثال‌هایی از این نوع داده‌ها می‌توان میزان بروز یک بیماری در مناطق مختلف، غلظت گازهای آلاینده هوا در میدان‌های مختلف یک شهر و درصد خلوص طلا در نقاط مختلف یک معدن را نام برد.

دو ویژگی مهم داده‌های فضایی، نمایش هر داده با موقعیت آن در فضای مورد مطالعه و وابستگی فضایی این داده‌هاست. معمولاً داده‌های فضایی که از موقعیت‌های مجاور جمع‌آوری می‌شوند، وابستگی بیش‌تری دارند و با افزایش فاصله بین موقعیت داده‌ها، این وابستگی کاهش می‌یابد. بدیهی است که به دلیل عدم برقراری شرط استقلال در داده‌های فضایی، تحلیل آماری این‌گونه داده‌ها با روش‌های آماری معمول مقدور نیست. از این رو، شاخه آمار فضایی برای تحلیل این‌گونه مشاهدات پدید آمده و در حال توسعه می‌باشد. در حقیقت آمار فضایی شاخه‌ای از علم آمار است که به تحلیل داده‌های فضایی می‌پردازد. این شاخه از آمار می‌کوشد بین مقادیر مختلف یک متغیر و فاصله و جهت قرارگیری آن‌ها نسبت به هم ارتباطی برقرار کند. این ارتباط

^۱Spatial data

فضایی، ساختار فضایی^۲ نام دارد. از جمله مباحث مهم در تحلیل داده‌های فضایی، بررسی و مطالعه ساختار وابستگی فضایی و برآورد آن با استفاده از مشاهدات و همچنین پیش‌گویی در موقعیت‌های فاقد مشاهده و تعیین خطای پیش‌گویی است.

اگرچه گسترش آمار فضایی اغلب توسط آماردانان صورت گرفته است، اما پیدایش و کاربرد عملی این شاخه از آمار صرفاً از طریق آماردانان انجام نشده است. توسعه این شاخه به وسیله زمین‌شناسان تحت عنوان زمین‌آمار^۳ صورت گرفته است. به همین سبب واژه‌پردازی و اصطلاحات آن با ادبیات آمار کلاسیک قدری متفاوت است. نخستین بار به دنبال روند تکاملی ذخایر معدنی که از قبل از سال ۱۹۶۰ آغاز شده بود، ماترون با انتشار مقاله بنیادین خود در سال ۱۹۶۲ توانست شاخه جدیدی از علم آمار تحت عنوان زمین‌آمار را پایه‌گذاری کند. زمین‌آمار به شاخه‌ای از علم آمار گفته می‌شود که خواص داده‌های فضایی را تحلیل می‌کند. در این شاخه، ناحیه جغرافیایی یک ناحیه چگال^۴ است، به این معنی که می‌توان بین هر دو موقعیت مکانی یک مکان دیگر را در نظر گرفت.

در شکل‌گیری آمار فضایی افراد مختلفی سهمیم بوده‌اند و به‌طور همزمان با ماترون فعالیت داشته‌اند. از جمله گاندین (۱۹۶۳) بحثی به نام تحلیل عینی^۵ را با مفاهیم بسیار نزدیک به تعاریف ماترون و دیگران در زمینه هواشناسی پیشنهاد کرد و از این طریق روش‌هایی برای تحلیل داده‌های فضایی ابداع نمود. ماترون (۱۹۶۳) پیش‌گویی فضایی کریگینگ^۶، که منتسب به دی.جی. کریگ است را مطرح نمود. کریگ یک مهندس معدن آفریقایی جنوبی بود که در معادن طلا کار می‌کرد و به تجربه فهمید که از اطلاعات بلوک‌های مجاور هم می‌توان برای پیش‌گویی عیار طلا استفاده کرد. دیوید (۱۹۷۷) از زمین‌آمار برای برآورد منابع ذخیره طلا استفاده کرد. جورنل و هیوجبرگ (۱۹۷۸) به طور مفصل روش‌های زمین‌آمار را با ادبیات خاص شرح دادند. ریپلی (۱۹۸۱) در کتابی تحت عنوان ”آمار فضایی“ پیش‌گویی فرآیندهای تصادفی را مطرح کرد. کرسی (۱۹۹۳) با انتشار کتاب خود با عنوان ”آمار برای داده‌های فضایی“، موجب پیشرفت گسترده این شاخه از آمار گردید.

داده‌های شمارشی همبسته فضایی در بسیاری از تحقیقات کاربردی مانند پزشکی، اقتصاد، زیست‌شناسی و علوم اجتماعی مورد توجه محققان قرار دارند. اولین روش گسترده تحلیل داده‌های فضایی استفاده از روش کریگینگ بود. کریگینگ یک روش درونیایی است که ویژگی کلیدی آن وابستگی به مکان می‌باشد. در این روش به مشاهدات نزدیک‌تر، به دلیل وابستگی بیشتر، وزن بزرگ‌تر و به مشاهدات دورتر وزن کوچک‌تری تخصیص داده می‌شود. در ۵۰ سال گذشته، از زمانی که کریگینگ معرفی شد، تعمیم‌هایی به‌منظور بهبود بخشیدن به

^۲ Spatial structure

^۳ Geostatistics

^۴ Dense

^۵ Objective analysis

^۶ Kriging

پیش‌گویی‌ها مطرح شده‌اند، که این تعمیم‌ها شامل کریگینگ نشانگر^۷ (سولو، ۱۹۸۶؛ کرسی، ۱۹۹۳)، کریگینگ معمولی^۸ (کرسی، ۱۹۸۸؛ جورنل و روسی، ۱۹۸۹)، کریگینگ عام^۹ (لیسوسکین و دیزایشکر، ۲۰۰۶؛ دیل زیمرمن و همکاران، ۱۹۹۹)، و هم کریگینگ^{۱۰} (گورث، ۱۹۹۸) می‌باشند. اگرچه کریگینگ سال‌های زیادی پذیره این که مدل زیربنایی داده‌ها مانا^{۱۱} (معمولا مانای مرتبه دوم) و یک فرآیند تصادفی گاوسی است را حفظ کرد، اما تحلیل داده‌های فضایی با این پذیره‌ها پیچیده‌تر و با مشکل مواجه شدند. در بسیاری از موقعیت‌های کاربردی با داده‌هایی مواجه می‌شویم که نرمال نیستند. یک نگرانی اصلی، داده‌های شمارشی فضایی و نسبت‌های فضایی بود که هر دو اساسا غیر نرمال هستند. به همین دلیل دو رده مدل‌های آمیخته خطی تعمیم‌یافته^{۱۲} (GLMMs) و مدل‌های سلسله‌مراتبی بیزی در سال ۱۹۹۰ ظهور کردند. مدل‌های آمیخته خطی تعمیم‌یافته به عنوان تعمیمی از مدل‌های خطی تعمیم‌یافته^{۱۳} (GLMs) برای توزیع‌های مختلف معرفی شدند. گات وی و استروپ (۱۹۷۷) دریافتند که مدل‌های آمیخته خطی تعمیم‌یافته فضایی^{۱۴} می‌توانند برای پیاده‌سازی ساختار همبستگی فضایی، با توزیع پاسخ غیرنرمال استفاده شوند. تاکنون مدل‌های گوناگونی برای تحلیل داده‌های فضایی غیرنرمال معرفی شده‌اند. نخستین مدل هم کریگینگ دوجمله‌ای، توسط لاجونی (۱۹۹۱) برای پیش‌گویی نرخ فضایی بر اساس نسبت نمونه‌ای مشاهده‌شده از یک بیماری نادر، پیشنهاد شد. الیور و همکاران (۱۹۹۸)، از این مدل برای تحلیل نرخ سرطان کودکان استفاده کردند. مونستیز و همکاران (۲۰۰۶) استفاده از مدل کریگینگ پواسون را برای داده‌های شمارشی همبسته فضایی پیشنهاد دادند.

در مواردی که ساختار ناهمگن موجود در داده‌های شمارشی یک جامعه متناهی فضایی موجب بیش‌پراکندگی می‌شود، عدم انعطاف توزیع دوجمله‌ای استفاده از آن را محدود می‌کند. برای مرتفع کردن این مساله از رهیافت جانشین آن مبتنی بر توزیع بتا-دوجمله‌ای استفاده می‌کنیم. انعطاف‌پذیری توزیع بتا-دوجمله‌ای باعث شده است که این توزیع در برازش مجموعه داده‌های بیش‌پراکنده سهم به‌سزایی را ایفا کند. در این پایان‌نامه، به پیش‌گویی فضایی با استفاده از توزیع بتا-دوجمله‌ای می‌پردازیم که انعطاف توزیع بتا-دوجمله‌ای موجب می‌شود تا به نتایج دقیق‌تری در استنباط برسیم. در ادامه، ابتدا به معرفی برخی از مفاهیم کلیدی آمار فضایی می‌پردازیم. بخش‌هایی از این فصل از کتاب محمدزاده (۱۳۸۹) با عنوان "آمار فضایی و کاربردهای آن"، برداشت شده است.

^۷Indicator kriging

^۸Ordinary kriging

^۹Universal kriging

^{۱۰}Cokriging

^{۱۱}Stationary

^{۱۲}Generalized linear mixed models

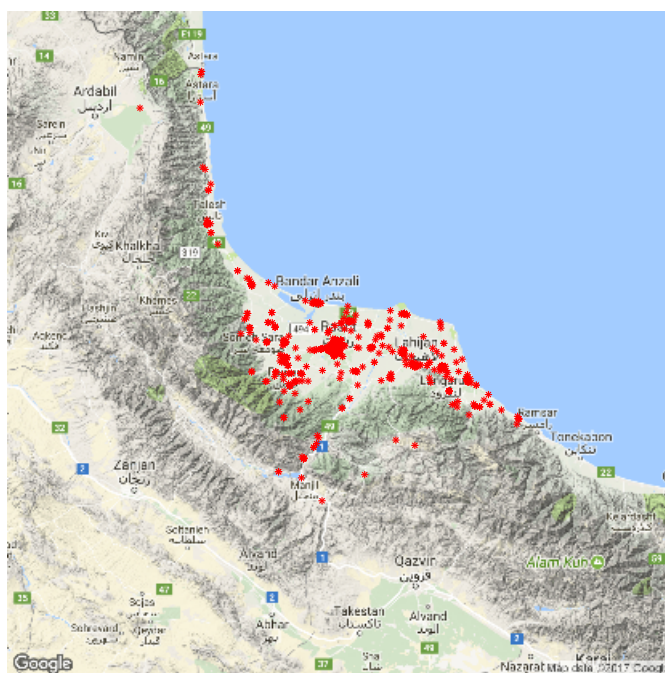
^{۱۳}Generalized linear models

^{۱۴}Spatial generalized linear models

۱.۱.۱ داده‌های فضایی

داده‌های فضایی مشاهداتی هستند که وابستگی آن‌ها ناشی از موقعیت و مکان قرار گرفتن آن‌ها در فضای (جغرافیایی) مورد مطالعه است. داده‌های فضایی بر اساس این که مکان‌های فضایی و متغیرهای تصادفی مرتبط به چه صورت باشند، به سه گروه داده‌های زمین‌آماري^{۱۵}، داده‌های شبکه‌ای^{۱۶} و الگوهای نقطه‌ای^{۱۷} تقسیم می‌شوند.

داده‌های زمین‌آماري: این نوع داده‌ها در موقعیت‌های ثابت و مشخص در ناحیه‌ای چگال مشاهده می‌شوند. یعنی بین دو موقعیت مفروض، ممکن است موقعیت‌های دیگری هم وجود داشته باشند. متغیر مورد بررسی ممکن است گسسته یا پیوسته باشد. به عنوان مثال می‌توان به مقدار ریزش باران ثبت شده در ایستگاه‌های هواشناسی، تعداد نوعی جانور دریایی در یک مجموعه از مکان‌های نمونه‌گیری شده و محل سکونت افراد مبتلا به سرطان معده در استان گیلان (شکل ۱.۱) اشاره کرد. هدف اصلی از تحلیل این گونه داده‌ها، پیش‌گویی پاسخ در مکان‌های مشاهده نشده است.



شکل ۱.۱: محل سکونت افراد مبتلا به سرطان معده در استان گیلان

داده‌های شبکه‌ای: این نوع داده‌ها مربوط به نواحی غیرچگال هستند، یعنی بین مکان‌های مشاهده شده مکان دیگری وجود ندارد. این مکان‌ها ممکن است منظم یا نامنظم باشند. تعداد افراد مبتلا به سرطان معده در شهرهای مختلف استان گیلان و تعداد حوادث رانندگی

^{۱۵} Geostatistical data

^{۱۶} Lattice data

^{۱۷} Point patterns

یا نرخ رشد اقتصادی هر استان مثال‌هایی از داده‌های شبکه‌ای هستند. در تحلیل داده‌های شبکه‌ای هدف عمده مدل‌بندی احتمالاتی مشاهدات است.

الگوهای نقطه‌ای: در این نوع داده‌ها موقعیت مشاهده‌شده، خود یک متغیر تصادفی است. این داده‌ها به سه دسته به‌طور کامل تصادفی فضایی^{۱۸} (CSR)، منظم^{۱۹} و خوشه‌ای^{۲۰} تقسیم و به مدل‌بندی آن‌ها اقدام می‌شود. به‌عنوان مثال می‌توان به موقعیت گونه‌ای از درختان در یک ناحیه جنگلی و محل تخلفات رانندگی در یک مسابقه رالی (شکل ۲.۱) اشاره کرد.



شکل ۲.۱: محل تخلفات رانندگی در یک مسابقه رالی

۲.۱.۱ مدل‌های آماری فضایی: میدان تصادفی

برای تحلیل داده‌های فضایی ابتدا باید یک مدل آماری در نظر گرفته شود. در آمار فضایی معمولاً مدل‌بندی داده‌ها با تعریف یک میدان تصادفی^{۲۱} صورت می‌گیرد. بنابراین ابتدا به معرفی میدان تصادفی و برخی از ویژگی‌های آن که در ادامه مورد استفاده قرار می‌گیرند، می‌پردازیم.

تعریف ۱.۱.۱. (میدان تصادفی). میدان تصادفی مجموعه‌ای از متغیرهای تصادفی مانند

$$Z(\cdot) = \{Z(s); s \in D \subseteq R^d, d \geq 1\}$$

است که در آن $z(s)$ متغیر تصادفی در اندیس s و D یک مجموعه اندیس‌گذار است.

تعریف ۲.۱.۱. (توابع میانگین، واریانس و کوواریانس میدان تصادفی). برای میدان تصادفی $Z(\cdot)$ میانگین در موقعیت s و کوواریانس در موقعیت‌های s_1 و s_2 به‌ترتیب به صورت

$$\mu(s) = E[Z(s)], s \in D$$

$$C(s_1, s_2) = Cov[Z(s_1), Z(s_2)] = E[(Z(s_1) - \mu(s_1))(Z(s_2) - \mu(s_2))], s_1, s_2 \in D$$

^{۱۸}Complete spatial randomness

^{۱۹}Regular

^{۲۰}Cluster

^{۲۱}Random field

تعریف می‌شوند. برای $s_1 = s_2 = s$ ، واریانس میدان تصادفی $Z(\cdot)$ در مکان s به صورت

$$Var[Z(s)] = E[Z(s) - \mu(s)]^2 = C(s, s)$$

حاصل می‌شود. هر میدان تصادفی $Z(\cdot)$ را می‌توان به صورت

$$Z(s) = \mu(s) + \delta(s), s \in D$$

تجزیه کرد که در آن $\mu(\cdot)$ تغییرات بزرگ مقیاس^{۲۲} یا روند و $\delta(\cdot)$ فرآیند خطا یا تغییرات کوچک مقیاس^{۲۳} میدان تصادفی نامیده می‌شوند. تغییرات کوچک مقیاس ممکن است از خطای اندازه‌گیری و تغییرات بزرگ مقیاس از تغییرپذیری بین موقعیت‌های مشاهده‌شده ناشی شده باشند.

به طور معمول، تحلیل داده‌های فضایی بر اساس مشاهدات نمونه دشوار است، اما در نظر گرفتن برخی ویژگی‌های میدان تصادفی موجب ساده‌تر شدن مساله خواهند شد که در ادامه به معرفی آن‌ها می‌پردازیم.

تعریف ۳.۱.۱. (مانای ذاتی). میدان تصادفی $Z(\cdot)$ را مانای ذاتی^{۲۴} گوئیم هرگاه

$$1. \text{ میانگین میدان تصادفی ثابت یا مستقل از } s \text{ باشد، یعنی } E(Z(s)) = \mu$$

۲. اختلاف واریانس در دو نقطه s_1 و s_2 فقط تابعی از فاصله موقعیت‌های s_1 و s_2 باشد. یعنی

$$Var(Z(s_1) - Z(s_2)) = 2\gamma(s_1 - s_2), s_1, s_2 \in D.$$

که در آن $\gamma(\cdot)$ تابعی از نوع فاصله است.

تعریف ۴.۱.۱. (مانای مرتبه دوم). میدان تصادفی $Z(\cdot)$ را مانای مرتبه دوم^{۲۵} گوئیم هرگاه علاوه بر مانایی در میانگین، یعنی $E(Z(s)) = \mu$ ، کوواریانس $Z(s_1)$ و $Z(s_2)$ فقط تابعی از فاصله موقعیت‌های s_1 و s_2 باشد، یعنی

$$Cov(Z(s_1), Z(s_2)) = C(s_1 - s_2), s_1, s_2 \in D \subset R^d.$$

بدیهی است که تحت مانایی مرتبه دوم

$$Var(Z(s)) = Cov(Z(s), Z(s)) = C(0) = \sigma^2$$

یعنی واریانس میدان تصادفی به موقعیت فضایی وابسته نیست و تغییرپذیری میدان در همه جا یکسان است.

^{۲۲} Large scale variation

^{۲۳} Small scale variation

^{۲۴} Intrinsic stationary

^{۲۵} Second order stationary

تعریف ۵.۱.۱. (مانای قوی). میدان تصادفی $Z(\cdot)$ مانای قوی نامیده می‌شود، هرگاه برای تمامی موقعیت‌های $s_1, \dots, s_n$ و $h \in R^d$ ، $(Z(s_1), \dots, Z(s_n)) \stackrel{D}{=} (Z(s_1 + h), \dots, Z(s_n + h))$ باشد. به بیان دیگر در صورت انتقال موقعیت‌های $s_1, \dots, s_n$ در راستای h ، توزیع توام تغییر نمی‌کند.

تعریف ۶.۱.۱. (میدان تصادفی گاوسی). میدان تصادفی $Z(\cdot)$ گاوسی نامیده می‌شود، هرگاه برای هر $m \geq 1$ توزیع توام $(Z(s_1), \dots, Z(s_m))$ نرمال چندمتغیره باشد.

قبل از ورود به بخش بعدی، با توجه به استفاده از نرم اقلیدسی در ادامه، ابتدا این تابع را تعریف می‌کنیم.

تعریف ۷.۱.۱. (تابع نرم اقلیدسی). فرض کنید $1 \leq p < \infty$ باشد، تابع p -نرم، برای بردار n بعدی $x = (x_1, x_2, \dots, x_n)$ از مقادیر حقیقی، به صورت زیر تعریف می‌شود:

$$\|x\|_p = \left(\sum_{i=1}^n |x_i|^p \right)^{\frac{1}{p}}.$$

اگر مقدار $p = 2$ باشد، نرم حاصل را اقلیدسی می‌نامند.

۳.۱.۱ ساختار همبستگی فضایی

در آمار فضایی، مشابه مفاهیم واریانس و کواریانس در آمار کلاسیک، ساختار همبستگی فضایی داده‌ها توسط تابع تغییرنگار^{۲۶} و هم‌تغییرنگار^{۲۷} تعیین و در تحلیل داده‌های فضایی مورد استفاده قرار می‌گیرد، که به ترتیب به تعریف آن‌ها می‌پردازیم.

تعریف ۸.۱.۱. (تغییرنگار). واریانس تفاضل بین مقادیر میدان تصادفی در دو موقعیت s و $s + h$ را تغییرنگار می‌نامند که به صورت

$$2\gamma(h) = \text{Var}(Z(s+h) - Z(s)) \quad (1.1)$$

تعریف می‌شود. تابع $\gamma(h)$ نیم‌تغییرنگار^{۲۸} نامیده می‌شود.

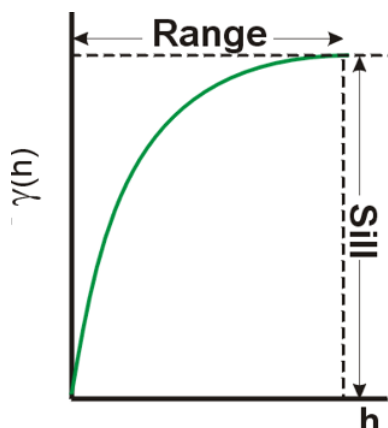
کوچک بودن تغییرنگار، نشانگر وابستگی زیاد و بزرگ بودن آن، نشانگر وابستگی کم میدان تصادفی است. تغییرنگار را می‌توان به عنوان معیاری برای نشان دادن تاثیرگذاری کمیت یک موقعیت بر کمیت موقعیت‌های دیگر در فضای مورد مطالعه به کار برد. نمودار تابع تغییرنگار در شکل ۳.۱ نشان داده شده است. اگر تغییرنگار برای تمام مقادیر h ثابت

^{۲۶} Variogram

^{۲۷} Covariogram

^{۲۸} Semi-variogram

باشد، داده‌ها همبستگی فضایی ندارند یا همبستگی بین آن‌ها بسیار ضعیف است. برعکس شیب نمودار تغییرنگار بیانگر همبستگی فضایی داده‌هاست. تغییرنگار سه پارامتر دامنه^{۲۹}، ازاره^{۳۰} و اثرقطعه‌ای^{۳۱} دارد که به ترتیب آن‌ها را تعریف می‌کنیم.



شکل ۳.۱: یکی از حالت‌های تابع تغییرنگار

تعریف ۹.۱.۱. (دامنه). فاصله‌ای که بعد از آن همبستگی داده‌های فضایی ناچیز و حتی مستقل از یکدیگر می‌شود، دامنه تغییرنگار نامیده می‌شود. تغییرنگار از این فاصله به بعد تغییر چندانی ندارد و به صورت تخت در می‌آید. از این رو داده‌های واقع در خارج از دامنه با روش‌های کلاسیک آماری قابل تحلیل هستند.

تعریف ۱۰.۱.۱. (ازاره). در صورتی که تغییرنگار به کران بالایی منتهی شود، آن کران ازاره تغییرنگار نامیده می‌شود و به صورت c در زیر تعریف می‌شود:

$$\lim_{||h|| \rightarrow \infty} \gamma(h) = c$$

تعریف ۱۱.۱.۱. (اثر قطعه‌ای). مقدار تغییرنگار در مبدا مختصات اثر قطعه‌ای نامیده می‌شود و به صورت c_0 در

$$\lim_{h \rightarrow 0} \gamma(h) = c_0$$

تعریف می‌شود. اثر قطعه‌ای معمولاً به دلیل خطاهای اندازه‌گیری یا تغییرات شدید کمیت مورد بررسی در موقعیت‌های نزدیک به هم ظاهر می‌شود. در رابطه با تغییرنگار می‌توان قضیه زیر را عنوان کرد.

^{۲۹}Range

^{۳۰}Sill

^{۳۱}Range

قضیه ۱.۱.۱. برای هر میدان تصادفی مانای ذاتی، تابع تغییرنگار همیشه منفی شرطی است، یعنی برای هر دنباله اعداد ثابت $\{a_i\}_{i=1}^m$ که $\sum_{i=1}^m a_i = 0$ ، $\sum_{i=1}^m \sum_{j=1}^m a_i a_j \gamma(s_i - s_j) \leq 0$ ، است (کرسی، ۱۹۹۳).

تعریف ۱۲.۱.۱. (هم تغییرنگار). کوواریانس دو متغیر $Z(s)$ و $Z(s+h)$ که به صورت

$$C(h) = \text{Cov}(Z(s), Z(s+h))$$

تعریف می شود را هم تغییرنگار گویند. اگر میدان تصادفی مانای مرتبه دوم باشد، آن گاه هم تغییرنگار به صورت $C(h) = E[Z(s)Z(s+h)] - \mu^2$ ساده می شود. علاوه بر این رابطه

$$\gamma(h) = C(0) - C(h)$$

بین تغییرنگار و هم تغییرنگار برقرار است، زیرا

$$\begin{aligned} 2\gamma(h) &= \text{Var}[Z(s+h) - Z(s)] \\ &= \text{Var}(Z(s+h)) + \text{Var}(Z(s)) - 2\text{Cov}[Z(s), Z(s+h)] \\ &= C(0) + C(0) - 2C(h). \end{aligned}$$

تابع هم تغییرنگار، تابعی معین مثبت است.

تعریف ۱۳.۱.۱. (همبستگی نگار). ضریب همبستگی بین $Z(s)$ و $Z(s+h)$ با شرط $C(0) > 0$ همبستگی نگار نام دارد و به صورت زیر تعریف می شود:

$$\rho(s, s+h) = \rho(h) = \frac{C(h)}{C(0)} = 1 - \frac{\gamma(h)}{C(0)}.$$

در حالت کلی توابع برداری تغییرنگار، کوواریانس و همبستگی نگار به اندازه و جهت بردار $h \in R^d$ بستگی دارند، اما اگر این سه فقط تابعی از $\|h\|$ ، یعنی نرم h ، در فضای R^d باشند و به جهت h بستگی نداشته باشند، آن ها را به ترتیب تغییرنگار، کوواریانس و همبستگی نگار همسان گرد^{۳۲} نامند و در غیر این صورت ناهمسان گرد نامیده می شوند.

تعریف ۱۴.۱.۱. (میدان تصادفی همسان گرد). میدان تصادفی $Z(\cdot)$ همسان گرد نامیده می شود، هرگاه توابع تغییرنگار، کوواریانس یا همبستگی نگار آن همسان گرد باشند.

مدل های معتبر تغییرنگار

توابعی که در شرط همیشه منفی شرطی صدق می کنند را مدل های معتبر تغییرنگار گویند. جورنل و هویچ برگتس (۱۹۷۸) مدل های پارامتری مختلفی را معرفی کرده اند که از آن جمله می توان به مدل های تغییرنگار گاوسی، نمایی، کروی، توانی، لگاریتمی، موجی، سهمی گون و ماترون اشاره کرد. در ادامه دو مدل گاوسی و کروی را معرفی می کنیم.

^{۳۲} Isotropic

۱. مدل گاوسی: تابع نیم‌تغییرنگار گاوسی به صورت

$$\gamma(h) = c_0 + c \left(1 - e^{-\frac{\|h\|}{a}}\right), h \in R^d, d \geq 1$$

تعریف می‌شود. به طوری که در آن پارامترهای a ، c_0 به ترتیب دامنه و اثر قطعه‌ای و $c + c_0$ نیز ازاره است.

۲. مدل کروی: تابع نیم‌تغییرنگار کروی به صورت

$$\gamma(h) = \begin{cases} c_0 + c \left(\frac{3}{4} \frac{\|h\|}{a} - \frac{1}{4} \frac{\|h\|^3}{a^3}\right) & 0 < \|h\| < a, h \in R^d, d = 1, 2, 3 \\ c + c_0 & \|h\| > a \end{cases}$$

می‌باشد، که دارای سه پارامتر a ، c و c_0 است.

۲.۱ پیش‌گویی فضایی

یکی از مسایل مهم در آمار فضایی، پیش‌گویی مقدار نامعلوم میدان تصادفی در موقعیت فاقد مشاهده s_0 بر اساس مشاهدات $Z = (Z(s_1), \dots, Z(s_n))'$ است. معمولاً با پذیرفتن مانایی میدان تصادفی، معلوم بودن ساختار همبستگی فضایی و در نظر گرفتن تابع زیان توان دوم خطا، پیش‌گوی بهینه از می‌نیم کردن میانگین شرطی $\hat{Z}(s_0) = E(Z(s_0) | Z)$ به دست می‌آید. بدون در نظر گرفتن پذیره‌های ساده‌کننده، محاسبه این میانگین شرطی از یک توزیع $(n+1)$ متغیره $(Z(s_0), Z')$ مقدور نیست. پذیره‌ای که معمولاً در نظر گرفته می‌شود، گاوسی بودن میدان تصادفی است. در این صورت توزیع توام $(Z(s_i), Z')$ نیز گاوسی است. همچنین پذیره گاوسی بودن میدان موجب می‌شود که پیش‌گوی $E(Z(s_0) | Z)$ نیز خطی شود و فقط به مقادیر $\mu(s_i) = E(Z(s_i))$ و $C(s_i, s_j) = Cov(Z(s_i), Z(s_j))$ وابسته باشد، که در آن $i, j = 1, \dots, n$. علاوه بر این، پذیره گاوسی بودن، هم‌واریانسی شرطی، یعنی عدم وابستگی $Var(Z(s_0) | Z)$ به Z ، را نیز فراهم می‌سازد. در این صورت تغییرات پیش‌گوی $E(Z(s_0) | Z)$ که با میانگین توان دوم خطای شرطی اندازه‌گیری می‌شود، از میانگین توان دوم خطای غیرشرطی، $E\left(Z(s_0) - \hat{Z}(s_0)\right)^2$ ، نامتمایز شده و محاسبه آن آسان می‌شود.

در آمار فضایی تعیین بهترین پیش‌گوی خطی ناریب برای $Z(s_0)$ کریگینگ نام دارد، که توسط ماترون (۱۹۶۳) به افتخار دی. جی. کریگ نام‌گذاری شده است. به عبارت دیگر کریگینگ یک روش پیش‌بینی در آمار فضایی به ویژه در زمین‌آمار می‌باشد، که معادل بهترین پیش‌گوی ناریب خطی است. در این روش برای پیش‌گویی در یک موقعیت مشخص، به مشاهدات نزدیک‌تر وزن بزرگ‌تر و به مشاهدات دورتر وزن کوچک‌تری داده می‌شود. همچنین وزن‌ها به گونه‌ای انتخاب می‌شوند که واریانس پیش‌گو کمینه شود. ویژگی دیگر کریگینگ این

است که در هر موقعیت فضایی واریانس پیش‌گو ارایه می‌شود و می‌توان توزیع خطاها را در فضای مورد بررسی به‌دست آورد.

فرض کنید پیش‌گویی $Z(s_0)$ بر اساس مشاهدات $Z(s_1), \dots, Z(s_n)$ در موقعیت‌های $s_1, \dots, s_n$ مورد نظر باشد. پیش‌گوی مکانی به‌صورت ترکیب خطی $\hat{Z}(s_0) = \sum_{i=1}^n \lambda_i Z(s_i) + k$ است، که در آن ضرایب λ_i و k طوری تعیین می‌شوند که $\hat{Z}(s_0)$ ناریب و دارای کم‌ترین واریانس در بین همه پیش‌گوهای خطی ناریب باشند. چنانچه میانگین میدان تصادفی $Z(s)$ ثابت، معلوم یا نامعلوم باشد، شرط ناریبی کریگینگ مستلزم آن خواهد بود که $k = 0$ و $\sum_{i=1}^n \lambda_i = 1$ باشند. بنابراین بسته به این‌که میانگین میدان تصادفی معلوم، ثابت نامعلوم یا تابعی از موقعیت‌ها باشد، کریگینگ به انواع کریگینگ ساده، معمولی و عام دسته‌بندی می‌شود، که در ادامه به معرفی آن‌ها می‌پردازیم.

۱.۲.۱ کریگینگ معمولی

کریگینگ معمولی توسط ماترون (۱۹۷۱)، تحت شرط‌های زیر معرفی شده است:

۱. میدان تصادفی به‌صورت

$$Z(s) = \mu + \delta(s) \quad (2.1)$$

است، که در آن $\mu \in \mathbb{R}$ مقداری ثابت اما نامعلوم و $\delta(s)$ میدان تصادفی مانای ذاتی با میانگین صفر است.

۲. پیش‌گو برای $Z(s_0)$ خطی و به‌صورت

$$\hat{Z}(s_0) = \sum_{i=1}^n \lambda_i Z(s_i)$$

با قید $\sum_{i=1}^n \lambda_i = 1$ است، که در آن‌ها وزن‌های کریگینگ و $Z(s_i)$ مشاهدات می‌باشند. شرط $\sum_{i=1}^n \lambda_i = 1$ ناریبی پیش‌گو را تضمین می‌کند، یعنی با استفاده از (۲.۱)

$$E[\hat{Z}(s_0)] = \mu$$

$$\sum \lambda_i E[Z(s_i)] = \mu \sum \lambda_i = \mu$$

بنابراین

$$\begin{aligned}
E \left[\hat{Z}(s_0) - Z(s_0) \right] &= E \left[\sum_{i=1}^n \lambda_i Z(s_i) - Z(s_0) \right] \\
&= E \left[\sum_{i=1}^n \lambda_i Z(s_i) \right] - E[Z(s_0)] \\
&= \sum_{i=1}^n \lambda_i E[Z(s_i)] - E(Z(s_0)) \\
&= \sum_{i=1}^n \lambda_i \mu - \mu \\
&= 0
\end{aligned}$$

از طرفی پیش‌گوی $Z(s_0)$ بهینه است هرگاه توان دوم خطای پیش‌گو، یعنی

$$\sigma_e^2 = E \left[Z(s_0) - \hat{Z}(s_0) \right]^2$$

در رده همه پیش‌گوهای خطی $\sum_{i=1}^n \lambda_i Z(s_i)$ ، با قید $\sum_{i=1}^n \lambda_i = 1$ می‌نیمم باشد.

فرض کنید در مدل (۲.۱)، $E[\delta(\cdot)] = 0$ و شرط مانایی مرتبه دوم برقرار و مدل دارای هم‌تغییرنگار معلوم $C(h)$ باشد. در این صورت توان دوم خطای پیش‌گو به صورت

$$\begin{aligned}
\sigma_e &= E \left[Z(s_0) - \hat{Z}(s_0) \right]^2 = E \left[Z(s_0) - \sum_{i=1}^n \lambda_i Z(s_i) \right]^2 \\
&= E \left[(Z(s_0) - \mu) - (\hat{Z}(s_0) - \mu) \right]^2 \\
&= E[Z(s_0) - \mu]^2 - 2E \left[(Z(s_0) - \mu) (\hat{Z}(s_0) - \mu) \right] + E[\hat{Z}(s_0) - \mu]^2 \\
&= C(0) - 2E \left[(Z(s_0) - \mu) \left(\sum_{i=1}^n \lambda_i Z(s_i) - \mu \right) \right] + \\
&\quad E \left[\left(\sum_{i=1}^n \lambda_i Z(s_i) - \mu \right) \left(\sum_{j=1}^n \lambda_j Z(s_j) - \mu \right) \right] \\
&= C(0) - 2 \sum_{i=1}^n \lambda_i E[(Z(s_0) - \mu)(Z(s_i) - \mu)] \\
&\quad + \sum_{i=1}^n \sum_{j=1}^n \lambda_i \lambda_j E[(Z(s_i) - \mu)(Z(s_j) - \mu)] \\
&= C(0) - 2 \sum_{i=1}^n \lambda_i C(s_0, s_i) + \sum_{i=1}^n \sum_{j=1}^n \lambda_i \lambda_j C(s_i, s_j)
\end{aligned}$$

حاصل می‌شود. حال با اضافه کردن ضریب لاگرانژ تابع هدف A را تشکیل می‌دهیم:

$$A = \sigma_e + 2\mu \left(\sum_{i=1}^n \lambda_i - 1 \right).$$

با استفاده از مشتق‌های جزئی خواهیم داشت

$$\begin{cases} \frac{\partial A}{\partial \lambda_i} = 0 - 2C(s_o, s_i) + 2 \sum_{j=1}^n \lambda_j C(s_i, s_j) + 2\mu = 0. \\ \frac{\partial A}{\partial \mu} = 0 \end{cases}$$

بنابراین می‌توان دستگاه زیر را تشکیل داد:

$$\begin{cases} \sum_{j=1}^n \lambda_j C(s_i, s_j) + \mu = C(s_o, s_i) \\ \sum_{i=1}^n \lambda_i = 1 \end{cases}$$

دستگاه معادلات بالا را می‌توان به صورت ماتریسی زیر نوشت:

$$\begin{bmatrix} C(s_1, s_1) & C(s_1, s_2) & \dots & C(s_1, s_n) & 1 \\ C(s_2, s_1) & C(s_2, s_2) & \dots & C(s_2, s_n) & 1 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ C(s_n, s_1) & C(s_n, s_2) & \dots & C(s_n, s_n) & 0 \end{bmatrix} \begin{bmatrix} \lambda_1 \\ \vdots \\ \lambda_n \\ \mu \end{bmatrix} = \begin{bmatrix} C(s_o, s_1) \\ C(s_o, s_2) \\ \vdots \\ C(s_o, s_n) \end{bmatrix}.$$

پیش‌گو و واریانس به ترتیب به صورت

$$\hat{Z}(s_o) = \lambda' Z$$

$$\sigma_k^2(s_o) = C(o) - \lambda' C + \mu$$

نتیجه می‌شوند، که در آن‌ها بردار $\lambda = (\lambda_1, \dots, \lambda_n)'$ و ضریب μ به صورت

$$\lambda' = \left[c + 1_n \frac{(1 - 1_n' \Sigma^{-1} c)}{1_n' \Sigma^{-1} 1_n} \right]' \Sigma^{-1}$$

$$\mu = -1 \frac{(1 - 1_n' \Sigma^{-1} c)}{1_n' \Sigma^{-1} 1_n}$$

محاسبه می‌شوند به طوری که

$$\Sigma = \begin{bmatrix} C(s_1, s_1) & C(s_1, s_2) & \dots & C(s_1, s_n) \\ C(s_2, s_1) & C(s_2, s_2) & \dots & C(s_2, s_n) \\ \vdots & \vdots & \ddots & \vdots \\ C(s_n, s_1) & C(s_n, s_2) & \dots & C(s_n, s_n) \end{bmatrix}$$

$$c = (C(s_o - s_1), \dots, C(s_o - s_n))$$

و 1_n برداری است که تمام n عنصر آن یک هستند. می‌دانیم که بین تغییرنگار و هم‌تغییرنگار رابطه

$$\gamma(h) = C(o) - C(h)$$

برقرار است. بنابراین دستگاه معادلات بر حسب تغییرنگار به صورت زیر قابل بازنویسی است:

$$\begin{bmatrix} \gamma(s_1, s_1) & \gamma(s_1, s_2) & \dots & \gamma(s_1, s_n) & 1 \\ \gamma(s_2, s_1) & \gamma(s_2, s_2) & \dots & \gamma(s_2, s_n) & 1 \\ \vdots & \vdots & \ddots & \vdots & \vdots \\ \gamma(s_n, s_1) & \gamma(s_n, s_2) & \dots & \gamma(s_n, s_n) & 0 \end{bmatrix} \begin{bmatrix} \lambda_1 \\ \vdots \\ \lambda_n \\ \mu \end{bmatrix} = \begin{bmatrix} \gamma(s_o, s_1) \\ \gamma(s_o, s_2) \\ \vdots \\ \gamma(s_o, s_n) \end{bmatrix}$$

واریانس پیش‌گو نیز به صورت زیر به دست می‌آید:

$$\sigma_{s_o}^2 = \sum_{i=1}^n \lambda_i \gamma(s_o - s_i) + \mu.$$

۲.۲.۱ کریگینگ ساده

کریگینگ ساده حالت خاصی از کریگینگ معمولی است، که در آن میانگین مقداری ثابت و معلوم است. بنابراین شرط ناریبی $\sum_{i=1}^n \lambda_i$ نیاز نیست. لازم به ذکر است که در دنیای واقعی میانگین نامعلوم است و کریگینگ ساده کاربردی ندارد.

۳.۲.۱ کریگینگ عام

کریگینگ عام، تعمیم‌یافته کریگینگ معمولی است که توسط ماترون (۱۹۶۹)، به عنوان یک پیش‌گوی ناریب، هنگامی که مقدار میانگین متغیر و نامعلوم است، معرفی شد.

۳.۱ مدل‌های خطی تعمیم‌یافته

مدلی که به طور معمول در رگرسیون مورد استفاده قرار می‌گیرد، مدل خطی^{۳۳} (LM) است. در یک مدل خطی با متغیر پاسخ Y و ماتریس طرح X ، تابع رگرسیونی، که همان میانگین شرطی $Y|X=x$ است، یک ترکیب خطی از X برای $E(Y|X=x)$ به صورت $x'\beta$ در نظر گرفته می‌شود. در این مدل، بردار β ، ضرایب رگرسیونی نامعلوم یا به عبارتی اثرات ثابت^{۳۴} نامیده می‌شوند. زمانی می‌توان از مدل‌های خطی استفاده کرد که تمام پذیره‌های آن برقرار باشد. یکی از این پذیره‌ها، توزیع نرمال برای متغیر پاسخ است. زمانی که این پذیره برای متغیر

^{۳۳} Liner model

^{۳۴} Fixed effects

پاسخ برقرار نباشد، نلدر و وردربرمن (۱۹۷۲) مدل‌های خطی تعمیم‌یافته را معرفی نمودند. مک‌کلا و نلدر (۱۹۸۹) مدل‌های خطی تعمیم‌یافته را برای مدل‌بندی متغیرهای پاسخ گسسته پیشنهاد دادند. در واقع این مدل‌ها تعمیمی از مدل‌های خطی بر اساس توزیع نرمال هستند که برای تحلیل پاسخ‌های عضو خانواده توزیع‌های نمایی گسترش یافته‌اند، که این تعمیم از دو قسمت تشکیل شده است. اول، توزیع‌هایی به‌جز توزیع نرمال در صورتی که متعلق به خانواده توزیع‌های نمایی باشند، می‌توانند توزیع پاسخ GLM را شامل شوند. دوم، میانگین پاسخ به‌طور مستقیم مدل‌سازی نمی‌شود بلکه با استفاده از تبدیلی از میانگین پاسخ به‌نام تابع پیوند^{۳۵} به متغیرهای تبیینی رگرسیونی مرتبط می‌شود.

تعریف ۱.۳.۱. (خانواده توزیع‌های نمایی). خانواده‌ای را عضو خانواده توزیع‌های نمایی متعارف برای متغیر تصادفی Y گویند، هرگاه بتوان تابع چگالی احتمال آن را به‌صورت زیر نوشت:

$$f_Y(y) = \exp \left\{ \frac{y\lambda - \psi(\lambda)}{a(\phi)} - c(y, \phi) \right\} \quad (۳.۱)$$

به‌طوری که ϕ و λ پارامترهای خانواده توزیع و $c(\cdot)$ ، $\psi(\cdot)$ و $a(\cdot)$ توابعی معلوم هستند.

در یک خانواده توزیع‌های نمایی امید ریاضی و واریانس به‌صورت زیر خواهند بود:

$$E(Y) = \frac{\partial \psi(\lambda)}{\partial \lambda} = \mu \quad \text{Var}(Y) = a(\phi) \frac{\partial^2 \psi(\lambda)}{\partial (\lambda)^2}. \quad (۴.۱)$$

از جمله توزیع‌های نمایی معروف می‌توان به توزیع‌های نرمال، دوجمله‌ای، پواسون، نمایی، گاما، بتا و گاوسی معکوس اشاره کرد. در یک مدل خطی تعمیم‌یافته، η ترکیبی خطی از متغیرهای تبیینی و پارامترهای رگرسیونی، به‌نام پیش‌گوی خطی^{۳۶} بوده به‌طوری

$$g(\mu) = \eta = x'\beta$$

$g(\cdot)$ یک تابع پیوند است. مدل‌های خطی تعمیم‌یافته دارای سه مولفه هستند:

۱. **مولفه تصادفی:** توزیع شرطی متغیر پاسخ را مشخص می‌کند. توزیع پاسخ، معمولاً از خانواده توزیع‌های نمایی مانند نرمال، دوجمله‌ای، بتا و پواسون انتخاب می‌شود، که انتخاب مولفه تصادفی بنا به ساختار متغیر پاسخ متفاوت است. برای مثال، اگر ساختار متغیر پاسخ به‌صورت متغیر شمارشی متناهی باشد، می‌توان توزیع دوجمله‌ای را به‌عنوان توزیع مولفه تصادفی انتخاب کرد.

۲. **مولفه منظم:** این مولفه نقش متغیرهای تبیینی را در قالب یک ترکیب خطی از متغیرهای تبیینی و پارامترهای رگرسیونی، η با نام پیش‌گوی خطی مشخص می‌کند. در واقع

^{۳۵} Link function

^{۳۶} Linear predictor

پیش‌گوی خطی به صورت

$$\eta = x' \beta$$

تعریف می‌شود.

۳. **تابع پیوند:** تابع پیوند به منظور ارتباط دادن مولفه منظم به مولفه تصادفی استفاده می‌شود. در واقع داریم

$$g(\mu) = \eta$$

یا به طور معادل

$$g^{-1}(\eta) = \mu.$$

تابع پیوند انواع مختلفی دارد. همان‌طور که عنوان شد موقعیت‌های کاربردی بسیاری وجود دارند که پذیره نرمال بودن متغیر پاسخ حتی به طور تقریبی برقرار نیست. به عنوان مثال، فرض کنید متغیر پاسخ یک متغیر گسسته شمارشی است. در مواردی که در بازه زمانی مشخصی حداکثر تعداد n مشخص نباشد از توزیع پواسون استفاده می‌کنیم که تابع پیوند آن لگاریتمی و به صورت زیر قابل تعریف است:

$$\eta = g(\mu) = \ln(\mu).$$

بنابراین

$$\mu = e^{\eta} > 0.$$

در بسیاری از موارد با شمارش بیمارانی که دارای یک بیماری خاص هستند، مواجه می‌شویم. در چنین مواردی که متغیر پاسخ شمارشی و حجم جامعه متناهی و مثلاً برابر n است، مدل معمول به صورت دوجمله‌ای تعریف می‌شود و از رگرسیون لجستیک برای مدل‌بندی داده‌ها استفاده می‌شود که تابع پیوند آن به صورت زیر می‌باشد:

$$\eta = \ln\left(\frac{\mu}{1-\mu}\right).$$

۱.۳.۱ مدل‌های آمیخته خطی تعمیم‌یافته

یکی از پذیره‌های اساسی مدل‌های خطی تعمیم‌یافته استقلال بین مشاهدات است. اما در مواردی مانند داده‌های فضایی که پاسخ‌ها به هم وابسته هستند، این پذیره برقرار نیست. در این موارد، تعمیمی از این مدل‌ها معروف به مدل‌های آمیخته خطی تعمیم‌یافته مورد استفاده قرار می‌گیرد. در واقع در این‌گونه موارد که بین مشاهدات همبستگی وجود دارد، معمولاً پذیره استقلال مشاهدات به استقلال شرطی تعدیل و همبستگی بین آن‌ها با اضافه کردن اثرات تصادفی از طریق متغیرهای پنهان به مدل لحاظ می‌شود. به عبارتی با افزودن اثرات

تصادفی به پیش‌گوی خطی در یک GLM، یک خانواده از مدل‌های منعطف که در موقعیت‌های عملی مختلفی کاربرد دارند، تشکیل می‌شود. متغیر پاسخ در GLMM مستقل شرطی بوده و دارای توزیعی از خانواده توزیع‌های نمایی با تابع چگالی احتمال (۳.۱) می‌باشد و به‌طور مشابه میانگین و واریانس آن از رابطه (۴.۱) به‌دست خواهند آمد. اگر میانگین شرطی برابر با $E(Y|X, Z, \alpha) = \mu$ باشد، آن‌گاه تابع پیوند روی این میانگین به‌عنوان ترکیبی خطی از هر دو عامل ثابت β و تصادفی α به مدل معرفی خواهد شد، یعنی

$$g(\mu) = X\beta + Z\alpha$$

که در آن Z ماتریس طرح متناظر با اثرات تصادفی α است.

۲.۳.۱ مدل‌های آمیخته خطی تعمیم‌یافته فضایی

باتوجه به عدم برقراری پذیره استقلال در بین مشاهدات غیرنرمال، از مدل‌های آمیخته خطی تعمیم‌یافته استفاده خواهیم کرد. در صورتی که همبستگی از نوع فضایی باشد، مدل‌های آمیخته خطی تعمیم‌یافته به مدل‌های آمیخته خطی تعمیم‌یافته فضایی تبدیل می‌شوند. برسلو و کلایتون (۱۹۹۳) از این مدل‌ها در مطالعات پزشکی استفاده کردند. دیگل و همکاران (۱۹۹۸) این مدل‌ها را برای تحلیل متغیرهای فضایی گسسته در یک جامعه ناحیه پیوسته به‌کار بردند. مدل مورد نظر در این پایان‌نامه یک مدل آمیخته خطی تعمیم‌یافته فضایی با پاسخ شمارشی در یک جامعه متناهی است.

۳.۳.۱ رگرسیون دوجمله‌ای

در برخی از موقعیت‌ها، وضعیت‌هایی وجود دارند که در آن متغیر پاسخ Y تنها دو برآمد دارد. مثلاً فشار خون بالا یا فشار خون پایین، یا پاسخ دادن یا پاسخ ندادن یک بیمار به یک دارو. در این قبیل موارد برآمد Y را می‌توان به‌صورت صفر یا یک کدگذاری کرد. در واقع می‌خواهیم برآمد یا احتمال آن را برای یک یا چند مقدار از متغیرهای تبیینی X پیش‌بینی کنیم. در این موارد از رگرسیون لجستیک استفاده می‌کنیم. به‌طور کلی رگرسیون لجستیک با رگرسیون خطی دو تفاوت دارد:

۱. در مدل رگرسیون خطی، کمیت کلیدی $E(Y|X = x)$ توسط یک رابطه خطی از X یا توسط تبدیلات این متغیرها بیان می‌شود. مانند $E(Y|X = x) = \beta_0 + \beta_1 x$. در این صورت مقادیر $E(Y|X = x)$ می‌تواند در بازه $-\infty$ تا $+\infty$ باشند. در مقابل، وقتی مقادیر پاسخ از نوع کیفی (اسمی و دودویی) باشند، میانگین که به معنای نسبت است، می‌تواند برآوردی برای $E(Y|X = x)$ ایجاد کند، لذا $0 \leq E(Y|X) \leq 1$. بنابراین

$$E(Y|X = x) = P(Y = 1|X = x) \times 1 + P(Y = 0|X = x) \times 0.$$

الگوی که بین صفر و یک محدود است و به‌طور مجانبی به صفر و یک میل می‌کند، توسط تابع پیوند لجیت به‌دست می‌آید و به‌صورت زیر بیان می‌شود:

$$p(x) = E(Y|X=x) = \frac{e^{\beta_0 + \beta_1 x}}{1 + e^{\beta_0 + \beta_1 x}}.$$

این الگو را با تبدیل ساده زیر می‌توان به الگوی خطی تبدیل کرد:

$$\text{logit}(p) = \ln\left(\frac{p(x)}{1-p(x)}\right) = \beta_0 + \beta_1 x.$$

۲. دومین اختلاف مربوط به توزیع شرطی متغیر پاسخ است. مدل رگرسیون خطی به‌صورت $y = x'\beta + \epsilon$ نوشته می‌شود که در آن $\epsilon \sim N(0, \sigma^2)$ در نظر گرفته می‌شود. اما این مطلب در مورد پاسخ‌های دودویی برقرار نیست. در این وضعیت مقدار متغیر پاسخ را با معلوم بودن x ، به‌صورت $y = \pi(x) + \epsilon$ در نظر می‌گیریم. اگر $y = 0$ باشد، آن‌گاه $\epsilon = -\pi(x)$ با احتمال $1 - \pi(x)$ است. بنابراین می‌توان گفت ϵ توزیعی با میانگین صفر و واریانس $\pi(x)[1 - \pi(x)]$ دارد، یعنی توزیع خطا به جای نرمال دوجمله‌ای است. بنابراین توزیع شرطی متغیر پاسخ از یک توزیع دوجمله‌ای با احتمال $\pi(x)$ پیروی می‌کند. پارامترهای β_0 و β_1 نیز معمولاً با روش ماکسیمم درست‌نمایی برآورد می‌شوند.

مدل رگرسیون لجستیک را می‌توان برای بیش از یک متغیر تبیینی نیز تعمیم داد. اگر متغیر پاسخ از نوع دودویی باشد و $x = (x_1, x_2, \dots, x_n)'$ متغیرهای تبیینی را تشکیل دهند، آن‌گاه مدل رگرسیون لجستیک به‌صورت زیر نوشته می‌شود:

$$P(Y=1|X=x) = E(Y|X=x) = \frac{\exp(\beta_0 + \beta_1 x_1 + \dots + \beta_n x_n)}{1 + \exp(\beta_0 + \beta_1 x_1 + \dots + \beta_n x_n)}$$

به‌طوری که

$$\ln\left(\frac{p}{1-p}\right) = \beta_0 + \beta_1 x_1 + \dots + \beta_n x_n.$$

۴.۳.۱ رگرسیون پواسون

رگرسیون پواسون مدلی بسیار پرکاربرد برای تحلیل داده‌های شمارشی است. مدل رگرسیون پواسون نوعی مدل خطی تعمیم یافته است که در آن متغیر پاسخ شمارشی بوده و از توزیع پواسون پیروی می‌کند. زمینه‌های کاربردی بسیاری برای رگرسیون پواسون وجود دارد. این کاربردها وقتی پیش می‌آیند که پاسخ طبیعی در مساله تنها مقادیر صحیح نامنفی $0, 1, 2, 3, \dots$ باشند. به‌عنوان مثال‌هایی از رگرسیون پواسون می‌توان به تعداد دفعات اهدای خون یا تعداد دفعات ترک اعتیاد اشاره کرد.

توزیع احتمال پواسون برای متغیر تصادفی Y به‌صورت زیر است:

$$f(y, \mu) = \frac{\mu^y e^{-\mu}}{y!} \quad y = 0, 1, 2, \dots$$

در مدل رگرسیون پواسون تابع پیوند طوری انتخاب می‌شود که معکوس آن یک تابع مثبت باشد. بنابراین تابع پیوند لگاریتمی یک انتخاب طبیعی است. یعنی

$$\ln(\mu) = x'\beta$$

که در آن $\mu = (x'\beta)$ ، β و x به ترتیب بردارهای k بعدی از پارامترها و متغیرهای تبیینی هستند. لگاریتم تابع درستنمایی این مدل برای n مشاهده به صورت زیر است:

$$\log L(\beta) = \sum_{i=1}^n \{-\exp(x'_i\beta) + y_i x'_i\beta - \log(y_i!)\}.$$

محاسبه برآورد پارامترها مستلزم حل معادلات غیرخطی $\sum_{i=1}^n \{y_i - \exp(x'_i\hat{\beta})\}x_i = 0$ توسط روش کم‌ترین توان‌های دوم تکراری وزنی^{۳۷} (IWLS) است. می‌توان نشان داد که برآورد پارامتر β توسط معادله $(X'WX)^{-1}(X'Wy^*)$ به صورت روش تکراری IWLS و به صورت عددی محاسبه می‌شود که در آن ماتریس وزنی W قطری با اعضای

$$w_{ij} = \frac{1}{\text{var}(y_i)} \left(\frac{\partial \mu_i}{\partial \eta_i} \right)^2$$

و $y^* = x'\hat{\beta} + (y - \mu) \frac{\partial \eta}{\partial \mu}$ بردار پاسخ اصلاح شده هستند. برآورد پارامترها، تحت شرایط نظم، به طور مجانبی دارای توزیع نرمال با میانگین β و ماتریس واریانس A^{-1} است. به عبارت دیگر $\sqrt{n}(\hat{\beta}_{ML} - \beta) \rightarrow N(0, A^{-1})$ که در آن

$$A = \lim_{n \rightarrow \infty} E \left[\frac{1}{n} \sum_i \exp(x'_i\beta) x_i x'_i \right].$$

علاوه بر این ویژگی سازگار بودن برآوردهای پارامترها نیز برقرار است. در مسائل کاربردی ممکن است با داده‌های شمارشی روبرو شویم که برآورد واریانس مشاهده شده از واریانس مورد انتظار بزرگ‌تر باشد. این ویژگی، بیش پراکنش^{۳۸} نامیده می‌شود. مدل پواسون برای برازش داده‌های شمارشی بیش پراکنش مناسب نیست (حسن زاده و کاظمی، ۱۳۹۰).

۴.۱ کریگینگ پواسون

مونستیز و همکاران (۲۰۰۶) مدل کریگینگ پواسون را برای مدل سازی مشاهدات شمارشی همبسته فضایی معرفی کردند. در ادامه به معرفی مدل و معادلات این نوع کریگینگ می‌پردازیم.

^{۳۷} Iterated weighted least square

^{۳۸} Overdispersion

۱.۴.۱ مشخصات مدل

میدان تصادفی $\{Z = Z(s); s \in D \subset \mathbb{R}^2\}$ را در نظر بگیرید. فرض کنید

$$Z(s) | Y(s) \sim \text{pois}(t(s)Y(s))$$

به‌طوری که در آن Z داده شمارشی مشاهده‌شده در مکان s ، $Y(s)$ میدان فضایی پنهان تولید کننده پارامتر نرخ پواسون و $t(s)$ نیز زمان مشاهدات در مکان s است. علاوه بر این، فرض کنید میدان تصادفی $Y(s)$ مثبت و مانای مرتبه دوم با میانگین m و واریانس σ_Y^2 باشد. تابع کوواریانس نیز فقط تابعی از فاصله مشاهدات به‌صورت $C_Y(\|s - s'\|)$ است. همچنین تابع کوواریانس

$$C_Y(\|s - s'\|) = \text{Cov}[Y(s), Y(s')]$$

می‌تواند در رابطه تغییرنگار یعنی، $\frac{1}{2} E[(Y(s) - Y(s'))^2] = \gamma_Y(s - s')$ بازچینی شود. پذیره دیگری بر روی میدان تصادفی پنهان به جز این که $Y \geq 0$ است وجود ندارد. در ادامه، به منظور سادگی $Z(s)$ ، $Y(s)$ و $t(s)$ را به‌ترتیب با Z_s ، Y_s و t_s نمایش می‌دهیم. همچنین در معادلات کریگینگ $C_{ss'}$ نشان‌دهنده تابع کوواریانس $C_Y(s - s')$ است. بنابراین با توجه به تعاریف بالا می‌توان نوشت

$$\begin{aligned} E[Z_s | Y_s] &= t_s Y_s, & E[Z_s] &= t_s E[Y_s] = m t_s \\ \text{Var}[Z_s | Y_s] &= t_s Y_s, & E[(Z_s)^2 | Y_s] &= t_s Y_s + (t_s Y_s)^2 \\ \text{Var}[Z_s] &= t_s^2 \sigma_Y^2 + m t_s. \end{aligned}$$

کوواریانس فضایی در دو مکان s_i و s_j نیز برابر است با

$$\begin{aligned} E[Z_s Z_{s'} | Y] &= \text{Cov}[Z_s, Z_{s'} | Y] + E[Z_s | Y_s] E[Z_{s'} | Y_{s'}] \\ &= \delta_{ss'} t_s Y_s + t_s t_{s'} Y_s Y_{s'} \end{aligned}$$

به‌طوری که در آن δ_{ij} دلتای کرونکر است، یعنی

$$\delta_{ij} = \begin{cases} 1 & i = j \\ 0 & i \neq j \end{cases}$$

امید ریاضی شرطی و کناری $\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}}$ به‌صورت

$$\begin{aligned} E\left[\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} | Y\right] &= \frac{1}{t_s} E[Z_s | Y_s] - \frac{1}{t_{s'}} E[Z_{s'} | Y_{s'}] \\ &= Y_s - Y_{s'} \end{aligned}$$

$$E\left[\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}}\right] = m - m = 0$$

هستند. امید ریاضی شرطی و کناری اختلاف توان دوم نسبت‌ها نیز به‌ترتیب به‌صورت

$$\begin{aligned} E \left[\left(\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} \right)^2 | Y \right] &= \frac{1}{t_s^2} E \left[(Z_s)^2 | Y_s \right] - \frac{1}{t_s t_{s'}} E \left[Z_s Z_{s'} | Y \right] - \frac{1}{t_{s'}^2} E \left[(Z_{s'})^2 | Y_{s'} \right] \\ &= \frac{Y_s}{t_s} - \frac{Y_{s'}}{t_{s'}} - \frac{1}{t_s t_{s'}} E \left[Z_s Z_{s'} | Y \right] + (Y_s - Y_{s'})^2 \end{aligned}$$

۹

$$\begin{aligned} E \left[\left(\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} \right)^2 \right] &= E \left[\frac{Y_s}{t_s} + \frac{Y_{s'}}{t_{s'}} - \frac{1}{t_s t_{s'}} E \left[Z_s Z_{s'} | Y \right] + (Y_s - Y_{s'})^2 \right] \\ &= m \left(\frac{t_s + t_{s'}}{t_s t_{s'}} \right) - \frac{1}{t_s t_{s'}} E \left[Z_s Z_{s'} \right] + \frac{1}{t_s t_{s'}} E \left[Z_s Z_{s'} \right] + \gamma_Y (Y_s - Y_{s'})^2 \end{aligned}$$

نتیجه می‌شوند. بنابراین برای واریانس می‌توان نوشت

$$Var \left[\left(\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} \right) \right] = E \left[\left(\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} \right)^2 \right] - E \left[\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} \right]^2.$$

بر اساس تعریف نیم‌تغییرنگار، خواهیم داشت

$$\frac{1}{2} Var \left[\left(\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} \right) \right] = \frac{m}{2} \left(\frac{t_s + t_{s'}}{t_s t_{s'}} \right) - \frac{1}{2} \delta_{ss'} \frac{m}{t_s} + \gamma_Y (s - s'). \quad (۵.۱)$$

برای $s \neq s'$ داریم

$$\gamma_Y (s - s') = \gamma_Z (s - s') - \frac{m}{2} \left(\frac{t_s + t_{s'}}{t_s t_{s'}} \right). \quad (۶.۱)$$

اگر $s = s'$ آن‌گاه معادله (۵.۱) به‌صورت زیر تبدیل می‌شود:

$$\gamma_Y (\circ) = \gamma_Z (\circ) = \circ. \quad (۷.۱)$$

علاوه بر این

$$\begin{aligned} Var \left[\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} | Y \right] &= E \left[\left(\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} \right)^2 | Y \right] - E^2 \left[\left(\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} \right) | Y \right] \\ &= \frac{Y_s}{t_s} - \frac{Y_{s'}}{t_{s'}} + (Y_s - Y_{s'})^2 - (Y_s - Y_{s'})^2 \\ &= \frac{Y_s}{t_s} + \frac{Y_{s'}}{t_{s'}} \end{aligned}$$

$$\begin{aligned} E \left[Var \left[\frac{Z_s}{t_s} - \frac{Z_{s'}}{t_{s'}} | Y \right] \right] &= E \left[\frac{Y_s}{t_s} + \frac{Y_{s'}}{t_{s'}} \right] \\ &= m \left(\frac{t_s + t_{s'}}{t_s t_{s'}} \right). \end{aligned}$$

۲.۴.۱ تغییرنگار کریگینگ پواسون

حال n مشاهده Z_i ، $i = 1, \dots, n$ ، با اندازه نمونه n از $Z(s_i)$ در زمان مشاهده شده t_i را در نظر بگیرید. برآورد تعدیل یافته برای تغییرنگار به صورت زیر به دست می‌آید:

$$\gamma_Y^*(h) = \frac{1}{N(h)} \sum_{i=1}^n \sum_{j=1}^n \frac{t_i t_j}{t_i + t_j} \left[\frac{1}{2} \left(\frac{Z_i}{t_i} - \frac{Z_j}{t_j} \right)^2 - \frac{m^*}{2} \left(\frac{t_i + t_j}{t_i t_j} \right) \right] I_d(i, j) \sim h$$

که در آن منظور از $I_d(i, j) \sim h$ تابع نشانگر برای نقاط s_i و s_j ای است که در فاصله h از هم قرار دارند. عدد $N(h)$ نیز تعداد نقاط همسایه در فاصله h و ثابت نرمال ساز است که برای همگن کردن تمام مکان‌های فضایی به صورت زیر تعریف می‌شود:

$$N(h) = \sum_{i,j} \frac{t_i t_j}{t_i + t_j} I_d(i, j) \sim h.$$

در معادله تغییرنگار m^* نیز برآورد میانگین Y است. عبارت تصحیح اریبی، $-\frac{m^*}{2} \left(\frac{t_i + t_j}{t_i t_j} \right)$ ، به طور مستقیم از رابطه (۶.۱) به دست می‌آید. حالت ساده شده تغییرنگار به صورت زیر حاصل می‌شود:

$$\gamma_Y^*(h) = \frac{1}{2N(h)} \times \sum_{ij} \left(\frac{t_i t_j}{t_i + t_j} \left(\frac{Z_i}{t_i} - \frac{Z_j}{t_j} \right)^2 - m^* \right) I_d(i, j) \sim h.$$

۳.۴.۱ معادلات کریگینگ پواسون

برای هر مکان فضایی جدید s_0 مقدار پیش گویی Y_0 می‌تواند به صورت ترکیب خطی از داده‌های مشاهده شده در زمان t_α با وزن‌های λ_i ، $i = 1, 2, \dots, m$ ، نوشته شود. یعنی

$$Y_0^* = \sum_{i=1}^n \lambda_i \frac{Z_i}{t_i}.$$

بهترین وزن‌های λ_i باید به گونه‌ای انتخاب شوند که مقادیر پیش گویی در مکان جدید نااریب باشند و کمترین واریانس پیش گو را داشته باشند. بنابراین با در نظر گرفتن شرط نااریبی Y_0^* ، امید ریاضی شرطی و حاشیه‌ای در مکان جدید عبارتند از

$$\begin{aligned} E[Y_0^* | Y] &= \sum_{i=1}^n \frac{\lambda_i}{t_i} t_i Y_i = \sum_{i=1}^n \lambda_i Y_i \\ E[Y_0^*] &= \sum_{i=1}^n \lambda_i E[Y_i] = m \sum_{i=1}^n \lambda_i \\ \sum_{i=1}^m \lambda_i &= 1. \end{aligned}$$

امید ریاضی شرطی وکناری توان دوم خطای پیش‌گو به‌صورت

$$\begin{aligned}
 E \left[(Y_o^* - Y_o)^2 | Y \right] &= E \left[\left(\sum_{i=1}^n \lambda_i \frac{Z_i}{t_i} - Y_o \right)^2 | Y \right] \\
 &= \sum_{i=1}^n \sum_{j=1}^n \frac{\lambda_i}{t_i} \frac{\lambda_j}{t_j} E [Z_i Z_j | Y] + Y_o^2 - 2 Y_o \sum_{i=1}^n \frac{\lambda_i}{t_i} E [Z_i | Y] \\
 &= \sum_{i=1}^n \sum_{j=1}^n \lambda_i \lambda_j Y_i Y_j + \sum_{i=1}^n \frac{\lambda_i^2}{t_i} Y_i + Y_o^2 - 2 Y_o \sum_{i=1}^n \lambda_i Y_i \\
 E \left[(Y_o^* - Y_o)^2 \right] &= E \left[\sum_{i=1}^n \sum_{j=1}^n \lambda_i \lambda_j Y_i Y_j + \sum_{i=1}^n \frac{\lambda_i^2}{t_i} Y_i + Y_o^2 - 2 Y_o \sum_{i=1}^n \lambda_i Y_i \right] \\
 &= \sum_{i=1}^n \sum_{j=1}^n \lambda_i \lambda_j C_{ij} + \sum_{i=1}^n \frac{\lambda_i^2}{t_i} m + \sigma_Y^2 - 2 \sum_{i=1}^n \lambda_i C_{io} + 2 m^2 - 2 m^2 \\
 &= \sigma_Y^2 + \sum_{i=1}^n \frac{\lambda_i^2}{t_i} m + \sum_{i=1}^n \sum_{j=1}^n \lambda_i \lambda_j C_{ij} - 2 \sum_{i=1}^n \lambda_i C_{io}
 \end{aligned}$$

است. بنابراین با توجه به قید ناریبی، خواهیم داشت

$$Var(Y_o^* - Y_o)^2 = E \left[(Y_o^* - Y_o)^2 \right].$$

با می‌نیمم کردن واریانس خطای پیش‌گو، معادله کریگینگ پواسون به‌صورت زیر حاصل می‌شود:

$$\begin{aligned}
 \sum_{j=1}^n \lambda_j C_{ij} + \lambda_i \frac{m}{t_i} + \mu &= C_{io} \quad i = 1, \dots, n \\
 \sum_{i=1}^n \lambda_i &= 1.
 \end{aligned}$$

همچنین شکل نهایی واریانس خطای پیش‌گو با توجه به واریانس کریگینگ معمولی به‌صورت

$$Var(Y_o^* - Y_o) = \sigma_Y^2 - \sum_{i=1}^n C_{ij} - \mu$$

نتیجه می‌شود.

در بخش‌هایی از این پایان‌نامه به منظور سادگی $Z(s_i)$ و $Y(s_i)$ به ترتیب با Z_i و Y_i نمایش داده شده‌اند.

۵.۱ توزیع‌ها

در این بخش، دو توزیع آماری دیگر را که در این پایان‌نامه از آن‌ها استفاده خواهیم کرد، معرفی می‌کنیم.

۱.۵.۱ توزیع دوجمله‌ای

هرگاه متغیر تصادفی X تعداد موفقیت‌ها در n بار تکرار مستقل یک آزمایش برنولی با احتمال موفقیت π در هر آزمایش باشد ($0 < \pi < 1$) آن‌گاه X را یک متغیر تصادفی دوجمله‌ای گویند و آن را با نماد $\text{bin}(n, \pi)$ نشان می‌دهند. تابع جرم احتمال توزیع دوجمله‌ای به صورت زیر است:

$$f_X(x) = \binom{n}{x} \pi^x (1 - \pi)^{n-x}, \quad x = 0, 1, \dots, n$$

در این توزیع $E(X) = n\pi$ و $\text{Var}(X) = n\pi(1 - \pi)$ است.

۲.۵.۱ توزیع بتا

اگر متغیر تصادفی X دارای تابع چگالی احتمال

$$f_{\alpha, \beta}(x) = \frac{1}{B(\alpha, \beta)} x^{\alpha-1} (1-x)^{\beta-1}; \quad 0 < x < 1, \quad \alpha, \beta > 0$$

باشد، آن‌گاه X را دارای توزیع بتا با پارامترهای α و β نامیده و آن را با نماد $X \sim \text{Beta}(\alpha, \beta)$ نمایش می‌دهند. همچنین

$$B(\alpha, \beta) = \frac{\Gamma(\alpha)\Gamma(\beta)}{\Gamma(\alpha + \beta)}$$

و

$$E(X) = \frac{\alpha}{\alpha + \beta}$$

$$\text{Var}(X) = \frac{\alpha\beta}{(\alpha + \beta)^2(\alpha + \beta + 1)}.$$

فصل ۲

مدل بتا-دوجمله‌ای فضایی

۱.۲ مقدمه

پاسخ‌های شمارشی همبسته فضایی در جوامع متناهی در بسیاری از تحقیقات کاربردی در زمینه‌هایی مانند پزشکی، اقتصاد، علوم زیست‌شناسی، علوم اجتماعی، بهداشت و فیزیک مورد توجه محققان قرار دارند. ماهیت پاسخ در بسیاری از موارد به‌صورت، نرخ، نسبت یا درصد است. به‌عنوان نمونه، ممکن است نسبت افراد مبتلا به سرطان معده، پدیده مورد علاقه باشد. در این موارد که حجم کل جامعه عدد مشخصی است، مدل متداول مبتنی بر توزیع دوجمله‌ای تعریف می‌شود. ویژگی مهم این توزیع، کوچک‌تر بودن واریانس از میانگین است. اما در موارد متعددی، این پذیره برقرار نیست و با مساله بیش‌پراکنش مواجه می‌شویم. بنابراین، در چنین شرایطی توزیع دوجمله‌ای نمی‌تواند ماهیت داده‌ها را به‌خوبی بیان کند و استفاده از آن برای مدل‌بندی داده‌ها منجر به استنباط‌های آماری نادرست در برآورد پارامترها، برآورد ساختار وابستگی و پیش‌گویی فضایی داده‌ها می‌شود. با این مقدمه، مدل بتا-دوجمله‌ای، به دلیل وجود دو پارامتر شکل در توزیع بتا، جانشینی بهتر برای مدل دوجمله‌ای است. این توزیع تمام حالت‌های بیش‌پراکنش، کم‌پراکنش و پراکنش متعادل را برای داده‌ها در نظر می‌گیرد. مبانی ایده مدل بتا-دوجمله‌ای توسط پیرسون مطرح شد و به صورت گسترده توسط اسکلام آن در آزمون هوش (هوین، ۱۹۷۹؛ لرد، ۱۹۶۵؛ ویلکاکس، ۱۹۸۱)، آزمایش سم‌شناسی^۱

^۱Toxicology

(ویلیامز، ۱۹۷۵)، همه‌گیرشناسی^۲ (گرفتس، ۱۹۷۳)، نمایش رسانه‌ها (گرین، ۱۹۷۰) و رفتار خریداران (مسی و همکاران، ۱۹۷۰) استفاده کرد.

ویلیامز (۱۹۷۵) داده‌های مربوط به تاثیر دارویی خاص بر روی حیوانات آزمایشگاهی را مورد بررسی قرار داد. در آزمایش وی متغیر پاسخ، تعداد توله‌هایی بود که در دوره شیرخوارگی زنده مانده‌اند. توزیع دوجمله‌ای با پارامتر ثابت π گزینه معمول برای این نوع داده‌ها بود، اما نکته قابل توجه در آزمایش او آن بود که نسبت بقای توله‌ها در یک زایمان، از یک زایمان به دیگری تغییر می‌کرد. به عبارت دیگر داده‌ها بیش‌پراکندگی داشتند و مدل دوجمله‌ای نمی‌توانست آن‌ها را به‌خوبی مدل‌بندی کند. ویلیامز از توزیع بتا-دوجمله‌ای به‌عنوان یک مدل انعطاف‌پذیر برای تغییرپذیری بین نمونه‌ها استفاده کرد. در مطالعه اثر زایمان روی مدل‌بندی دز پاسخ در تراریختگی، کوپر و همکارانش (۱۹۸۶) نتیجه گرفتند که عدم به‌کارگیری تاثیر زایمان می‌تواند باعث کم‌برآوردی واریانس مرتبط با برآوردهای پارامترها باشد، به‌طوری که استفاده از درستمایی دوجمله‌ای برای مدل‌بندی داده‌های تراریختگی^۳ به نظر معقول نیست. آن‌ها توزیع بتا-دوجمله‌ای را برای معرفی درجه تغییر همبستگی درون زایمان پیشنهاد کردند. پل (۱۹۸۲) در تحلیل نسبت‌های جنین تاثیرپذیر در آزمایش‌های تراریختگی دریافت که مدل بتا-دوجمله‌ای نسبت به مدل دوجمله‌ای بهتر عمل می‌کند. پک (۱۹۸۶) به این نتیجه رسید که مدل بتا-دوجمله‌ای نسبت به مدل‌های معادل مانند دوجمله‌ای همبسته کوپر و هاسمن (۱۹۷۸) بهتر است. تارون (۱۹۸۲) برای بسیاری از انواع غده‌ها مشاهده کرد که نرخ کنترل عمر غده‌ها بسیار تغییرپذیرتر از آن است که پذیره دوجمله‌ای برای آن‌ها مناسب باشد، و توزیع بتا-دوجمله‌ای را برای نرخ عمر غده‌ها برازش داد. شوکرز (۲۰۰۳) نیز از توزیع بتا-دوجمله‌ای در ارزیابی کارایی دستگاه تشخیص هویت زیست‌سنجی^۴ استفاده کرد.

در این فصل ابتدا به معرفی مدل بتا-دوجمله‌ای و خواص آن می‌پردازیم. در ادامه مدل کریگینگ بتا-دوجمله‌ای که شامل معادلات کریگینگ و واریانس پیش‌گو می‌باشد را بیان می‌کنیم.

۲.۲ توزیع بتا-دوجمله‌ای

فرض کنید متغیر تصادفی شرطی $X|\pi$ نشان‌دهنده تعداد موفقیت‌ها در n آزمایش برنولی با احتمال موفقیت π باشد. اگر متغیر تصادفی π دارای توزیع پیشین بتا با پارامترهای شکل α و β باشد، به‌طوری که احتمال موفقیت π برای هر رخداد متغیر است، آن‌گاه متغیر تصادفی X دارای توزیع کناری بتا-دوجمله‌ای با پارامترهای n ، α و β است که تابع احتمال آن به‌صورت زیر به‌دست می‌آید:

^۲ Epidemiology

^۳ Teratology

^۴ Biometrics

فرض کنید

$$X|\pi \sim \text{bin}(n, \pi)$$

به‌طوری که

$$\pi \sim B(\alpha, \beta).$$

بنابراین تابع احتمال توام به‌صورت زیر نتیجه می‌شود:

$$f(x, \pi) = f(\pi) \cdot f(x|\pi) = \frac{\binom{n}{x}}{B(\alpha, \beta)} \pi^{\alpha+x-1} (1-\pi)^{\beta+n-x-1}$$

که در آن

$$B(\alpha, \beta) = \frac{\Gamma(\alpha) \Gamma(\beta)}{\Gamma(\alpha + \beta)}.$$

با انتگرال‌گیری از توزیع توام نسبت به π ، داریم

$$f(x) = \int_0^1 \frac{\binom{n}{x}}{B(\alpha, \beta)} \pi^{\alpha+x-1} (1-\pi)^{\beta+n-x-1} d\pi.$$

با ضرب $B(\alpha + x, \beta + n - x)$ خواهیم داشت

$$f(x) = \frac{\binom{n}{x} B(\alpha + x, \beta + n - x)}{B(\alpha, \beta)} \int_0^1 \frac{1}{B(\alpha + x, \beta + n - x)} \pi^{\alpha+x-1} (1-\pi)^{\beta+n-x-1} d\pi.$$

با توجه به این که

$$\int_0^1 \frac{1}{B(\alpha + x, \beta + n - x)} \pi^{\alpha+x-1} (1-\pi)^{\beta+n-x-1} d\pi = 1$$

پس تابع احتمال کناری به‌صورت

$$f(x) = \binom{n}{x} \frac{B(\alpha + x, \beta + n - x)}{B(\alpha, \beta)}, \quad x = 0, 1, \dots, n$$

حاصل می‌شود که همان توزیع بتا-دوجمله‌ای است و آن را با نماد $X \sim \text{beta-bin}(n, \alpha, \beta)$ نشان می‌دهند.

با استفاده از خواص توزیع گاما می‌دانیم

$$\Gamma(\alpha + x) = \Gamma(\alpha) \alpha(\alpha + 1) \cdots (\alpha + x - 1).$$

بنابراین با توجه به ویژگی‌های ذکر شده، تابع احتمال بتا-دوجمله‌ای به‌صورت زیرقابل بازنویسی است:

$$f(x) = \binom{n}{x} \frac{\alpha(\alpha + 1) \cdots (\alpha + x - 1) \beta(\beta + 1) \cdots (\beta + n - x - 1)}{(\alpha + \beta)(\alpha + \beta + 1) \cdots (\alpha + \beta + n - 1)}.$$

۱.۲.۲ میانگین و واریانس توزیع بتا-دوجمله‌ای

میانگین و واریانس توزیع بتا-دوجمله‌ای را می‌توان به صورت زیر محاسبه کرد:

$$\begin{aligned} E(X) &= \sum_{x=0}^n x \binom{n}{x} \frac{B(\alpha+x, \beta+n-x)}{B(\alpha, \beta)} \\ &= \frac{xn!}{x!(n-x)!} \frac{B(\alpha+x, \beta+n-x)}{B(\alpha, \beta)} \\ &= \sum_{x=1}^n n \binom{n-1}{x-1} \frac{B(\alpha+x, \beta+n-x)}{B(\alpha, \beta)} \\ &= n \sum_{x=0}^{n-1} \binom{n-1}{x} \frac{B(\alpha+x+1, \beta+n-x-1)}{B(\alpha, \beta)} \end{aligned}$$

اکنون با ضرب $B(\alpha+1, \beta)$ می‌توان نوشت

$$E(X) = \frac{nB(\alpha+1, \beta)}{B(\alpha, \beta)} \sum_{x=0}^{n-1} \binom{n-1}{x} \frac{B(\alpha+x+1, \beta+n-x-1)}{B(\alpha+1, \beta)}.$$

از آن جایی که

$$\sum_{x=0}^{n-1} \binom{n-1}{x} \frac{B(\alpha+x+1, \beta+n-x-1)}{B(\alpha+1, \beta)} = 1$$

می‌توان نوشت

$$\begin{aligned} E(X) &= \frac{nB(\alpha+1, \beta)}{B(\alpha, \beta)} = \frac{n\Gamma(\alpha+1)\Gamma(\beta)}{\Gamma(\alpha+\beta+1)} \frac{\Gamma(\alpha+\beta)}{\Gamma(\alpha)\Gamma(\beta)} \\ &= n\alpha \frac{\Gamma(\alpha)\Gamma(\beta)}{(\alpha+\beta)\Gamma(\alpha+\beta)} \frac{\Gamma(\alpha+\beta)}{\Gamma(\alpha)\Gamma(\beta)}. \end{aligned}$$

در نتیجه

$$E(X) = \frac{n\alpha}{\alpha+\beta}.$$

برای محاسبه واریانس طبق تعریف داریم

$$Var(X) = E(X^2) - E^2(X). \quad (1.2)$$

از آن جا که

$$E(X^2) = E(X(X-1)) + E(X)$$

بنابراین

$$\begin{aligned} E(X(X-1)) &= \sum_{x=0}^n (x)(x-1) \frac{n(n-1)(n-2)}{x(x-1)(x-2)(n-x)!} \frac{B(\alpha+x, \beta+n-x)}{B(\alpha, \beta)} \\ &= n(n-1) \sum_{x=2}^n \binom{n-2}{x-2} \frac{B(\alpha+x, \beta+n-x)}{B(\alpha, \beta)} \end{aligned}$$

با ضرب $B(\alpha + 2, \beta)$ داریم

$$E(X(X-1)) = \frac{n(n-1)B(\alpha+2, \beta)}{B(\alpha, \beta)} \sum_{x=0}^{n-2} \binom{n-2}{x} \frac{B(\alpha+x+2, \beta+n-x-2)}{B(\alpha+2, \beta)}.$$

باتوجه به

$$\sum_{x=0}^{n-2} \binom{n-2}{x} \frac{B(\alpha+x+2, \beta+n-x-2)}{B(\alpha+2, \beta)} = 1$$

نتیجه می‌شود

$$E(X(X-1)) = \frac{n(n-1)B(\alpha+2, \beta)}{B(\alpha, \beta)}.$$

بنابراین

$$\begin{aligned} E(X^2) &= n(n-1) \frac{B(\alpha+2, \beta)}{B(\alpha, \beta)} + \frac{n\alpha}{\alpha+\beta} \\ &= \frac{\Gamma(\alpha+2)\Gamma(\beta)}{\Gamma(\alpha+\beta+2)} \frac{\Gamma(\alpha+\beta)}{\Gamma(\alpha)\Gamma(\beta)} \\ &= \frac{(\alpha+1)\alpha\Gamma(\alpha)\Gamma(\beta)}{(\alpha+\beta+1)(\alpha+\beta)\Gamma(\alpha+\beta)} \frac{\Gamma(\alpha+\beta)}{\Gamma(\alpha)\Gamma(\beta)} \\ &= \frac{n(n-1)\alpha(\alpha+1)}{(\alpha+\beta+1)(\alpha+\beta)} + \frac{n\alpha}{\alpha+\beta}. \end{aligned}$$

با جای گذاری در رابطه (۱.۲) خواهیم داشت

$$\begin{aligned} Var(X) &= \frac{n^2\alpha^2 + n^2\alpha + n\alpha\beta}{(\alpha+\beta)(\alpha+\beta+1)} - \frac{n^2\alpha^2}{(\alpha+\beta)^2} \\ &= \frac{n\alpha\beta(\alpha+\beta+n)}{(\alpha+\beta)^2(\alpha+\beta+1)} \end{aligned}$$

۲.۲.۲ توزیع جعبه و مهره برای مدل بتا-دوجمله‌ای

توزیع بتا-دوجمله‌ای را می‌توان به صورت یک مدل جعبه برای مقادیر حقیقی مثبت α و β بیان کرد.

فرض کنید جعبه‌ای شامل α گوی سفید و β گوی سیاه باشد. یک گوی را به تصادف از جعبه خارج می‌کنیم. اگر گوی سفید مشاهده شود، آن را همراه با گوی سفید دیگری به جعبه بر می‌گردانیم. برای گوی سیاه نیز به همین صورت عمل می‌کنیم. اگر این عمل را n بار تکرار کنیم، آن گاه احتمال مشاهده x گوی سفید دارای توزیع بتا-دوجمله‌ای با پارامترهای n, α و β می‌باشد که به صورت زیر به دست می‌آید:

$$\begin{aligned} P(X=x) &= \binom{n}{x} \frac{\alpha(\alpha+1)\cdots(\alpha+x-2)(\alpha+x-1)}{(\alpha+\beta)(\alpha+\beta+1)\cdots(\alpha+\beta+x-1)} \frac{\beta(\beta+1)\cdots(\beta+n-x-2)(\beta+n-x-1)}{(\alpha+\beta+x)\cdots(\alpha+\beta+n-1)} \\ &= \binom{n}{x} \frac{\Gamma(\alpha+x)\Gamma(\beta+n-x)\Gamma(\alpha+\beta)}{\Gamma(\alpha+\beta+n)\Gamma(\alpha)\Gamma(\beta)} \\ &= \binom{n}{x} \frac{B(\alpha+x, \beta+n-x)}{B(\alpha, \beta)}. \end{aligned}$$

لازم به ذکر است که اگر گوی‌ها را به‌صورت تصادفی و با جای‌گذاری از جعبه خارج کنیم و دوباره به جعبه برگردانیم، و این عمل را n بار تکرار کنیم، آن‌گاه احتمال مشاهده x گوی سفید دارای توزیع دوجمله‌ای با پارامترهای n و $\pi = \frac{\alpha}{\alpha+\beta}$ می‌باشد. همچنین اگر گوی‌ها را به‌صورت تصادفی و بدون جای‌گذاری خارج کنیم، آن‌گاه احتمال مشاهده x گوی سفید دارای توزیع فوق هندسی است.

۳.۲.۲ فرآیند فضایی بتا-دوجمله‌ای

فرض کنید میدان تصادفی $\{Z = Z(s); s \in D \subset \mathbb{R}^2\}$ که در m مکان s_i ، $i = 1, 2, \dots, m$ مشاهده شده است، دارای توزیع شرطی $Z_i|Y_i \sim \text{bin}(n_i, Y_i)$ باشد، به‌طوری که احتمال موفقیت Y_i دارای توزیع بتا با پارامترهای α و β و ساختار وابستگی Σ_Y است. در واقع در این حالت میدان تصادفی Y_i همبسته فضایی و مدل زیربنایی، فرآیند بتا-دوجمله‌ای محسوب می‌شود. به عبارت دیگر وابستگی فضایی فقط در سطح Y_i وجود دارد. علاوه بر این، فرض کنید میدان تصادفی Y_i مانای مرتبه دوم باشد. یعنی میانگین توزیع بتا مقدار ثابت π و تابع کوواریانس فقط تابعی از فاصله مشاهدات به‌صورت $C_Y(\|s_i - s_j\|)$ است. افزون بر این، تابع تغییرنگار مدل را با $\gamma_Y(\|s_i - s_j\|)$ نشان می‌دهیم. بنابراین می‌توان نوشت

$$E(Y_i) = \frac{\alpha}{\alpha + \beta} = \pi \quad \text{Var}(Y_i) = \frac{\alpha\beta}{(\alpha + \beta)^2(\alpha + \beta + 1)} = \sigma_Y^2.$$

با این ویژگی‌ها متغیر تصادفی Z_i دارای توزیع کناری بتا-دوجمله‌ای با تابع احتمال

$$P(Z_i = z_i) = \frac{B(\alpha + z_i, \beta + n_i - z_i)}{B(\alpha, \beta)} \binom{n_i}{z_i} \quad z_i = 0, \dots, m$$

می‌باشد. با استفاده از قاعده مکرر امید ریاضی و واریانس شرطی، امید و واریانس توزیع بتا-دوجمله‌ای به‌صورت زیر نتیجه می‌شوند:

$$E(Z_i) = E[E(Z_i|Y_i)] = n_i E(Y_i) = n_i \pi$$

$$\text{Var}(Z_i) = \text{Var}[E(Z_i|Y_i)] + E[\text{Var}(Z_i|Y_i)] = n_i \sigma_Y^2 (n_i - 1) + n_i \pi (1 - \pi).$$

این روابط تاثیر میدان تصادفی پنهان بتا را با وارد کردن پارامترهای شکل α و β در میانگین و واریانس مدل بتا-دوجمله‌ای نسبت به دوجمله‌ای نشان می‌دهند. میانگین و واریانس توزیع دوجمله‌ای به‌ترتیب $n_i \pi (1 - \pi)$ و $n_i \pi$ است. بنابراین عبارت $n_i \sigma_Y^2 (n_i - 1)$ در واریانس توزیع بتا-دوجمله‌ای به‌عنوان عامل افزایش واریانس توزیع دوجمله‌ای عمل کرده و مساله بیش‌برازش داده‌ها را لحاظ می‌کند.

برای محاسبه کوواریانس شرطی بین متغیرهای Z_i و Z_j می‌توان نوشت

$$\begin{aligned} E(Z_i Z_j | Y) &= \text{Cov}(Z_i, Z_j | Y) + E(Z_i | Y_i) E(Z_j | Y_j) \\ &= \delta_{ij} n_i y_i (1 - y_i) + n_i n_j y_i y_j \end{aligned}$$

که در آن δ_{ij} دلتای کرونکر است. یعنی

$$\delta_{ij} = \begin{cases} 1 & i = j \\ 0 & i \neq j \end{cases}$$

در این صورت $Cov(Z_i, Z_j|Y) = 0$ هرگاه Z_i و Y_i ناهمبسته باشد و همچنین اگر $i = j$ آن‌گاه

$$Cov(Z_i, Z_j|Y) = Var(Z_i|Y_i).$$

هدف پیش‌گویی فضایی پاسخ‌های نرخ، با استفاده از توزیع بتا-دوجمله‌ای است. بنابراین در این مرحله، مدل از داده‌های شمارشی مشاهده‌شده بتا-دوجمله‌ای به نسبت‌های نمونه‌ای $\frac{Z_i}{n_i}$ تغییر می‌یابد. در ادامه ابتدا روابطی را بین نسبت‌های نمونه‌ای در دو مکان فضایی s_i و s_j بیان می‌کنیم که در نهایت این روابط منجر به تغییرنگار نظری خواهند شد.

امید ریاضی شرطی اختلاف نسبت‌های نمونه‌ای شرطی $\frac{Z_i}{n_i} - \frac{Z_j}{n_j}$ به صورت

$$\begin{aligned} E\left[\frac{Z_i}{n_i} - \frac{Z_j}{n_j} | Y\right] &= \frac{1}{n_i} E(Z_i | Y_i) - \frac{1}{n_j} E(Z_j | Y_j) \\ &= Y_i - Y_j \end{aligned}$$

است و امید ریاضی کناری آن عبارت است از

$$\begin{aligned} E\left[\frac{Z_i}{n_i} - \frac{Z_j}{n_j}\right] &= E\left(E\left[\left(\frac{Z_i}{n_i} - \frac{Z_j}{n_j}\right) | Y\right]\right) \\ &= E(Y_i - Y_j) \\ &= \pi - \pi = 0. \end{aligned}$$

به‌طور کلی فرض می‌کنیم که میدان تصادفی پنهان بتا مانا و دارای میانگین ثابت است. بنابراین اختلاف امید ریاضی حاشیه‌ای در دو مکان s_i و s_j صفر خواهد شد.

حال امید ریاضی شرطی اختلاف توان دوم نسبت‌ها را به دست می‌آوریم. در این جا فرض

می‌کنیم که تمام مکان‌های فضایی مجزا هستند، یعنی $s_i \neq s_j$. بنابراین

$$\begin{aligned} E \left[\left(\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right)^2 | Y \right] &= E \left[\left(\frac{Z_i}{n_i} + \frac{Z_j}{n_j} - 2 \frac{Z_i Z_j}{n_i n_j} \right) | Y \right] \\ &= \frac{1}{n_i} E(Z_i | Y_i) + \frac{1}{n_j} E(Z_j | Y_j) - \frac{2}{n_i n_j} E(Z_i Z_j | Y) \\ &= \frac{1}{n_i} [n_i Y_i + n_i^2 Y_i^2 - n_i Y_i^2] + \frac{1}{n_j} [n_j Y_j + n_j^2 Y_j^2 - n_j Y_j^2] - \\ &\quad \frac{2}{n_i n_j} [n_i n_j Y_i Y_j] \\ &= \left[\frac{Y_i}{n_i} + Y_i^2 - \frac{Y_i^2}{n_i} \right] + \left[\frac{Y_j}{n_j} + Y_j^2 - \frac{Y_j^2}{n_j} \right] - 2 Y_i Y_j \\ &= \frac{Y_i}{n_i} - \frac{Y_i^2}{n_i} + \frac{Y_j}{n_j} - \frac{Y_j^2}{n_j} + (Y_i - Y_j)^2 \\ &= \frac{1}{n_i} Y_i (1 - Y_i) + \frac{1}{n_j} Y_j (1 - Y_j) + (Y_i - Y_j)^2. \end{aligned}$$

امید ریاضی کناری اختلاف توان دوم نسبت‌ها عبارت است از

$$\begin{aligned} E \left[\left(\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right)^2 \right] &= E \left[E \left(\left(\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right)^2 | Y \right) \right] \\ &= E \left[\frac{1}{n_i} Y_i (1 - Y_i) + \frac{1}{n_j} Y_j (1 - Y_j) + (Y_i - Y_j)^2 \right] \\ &= \frac{1}{n_i} \pi - \frac{1}{n_i} (\sigma_Y^2 + \pi^2) + \frac{1}{n_j} \pi - \frac{1}{n_j} (\sigma_Y^2 + \pi^2) + 2 \gamma_Y (s_i - s_j) \\ &= \pi \frac{n_i + n_j}{n_i n_j} - \frac{n_i + n_j}{n_i n_j} (\sigma_Y^2 + \pi^2) + 2 \gamma_Y (s_i - s_j) \\ &= \frac{n_i + n_j}{n_i n_j} (\pi - \sigma_Y^2 - \pi^2) + 2 \gamma_Y (s_i - s_j). \end{aligned}$$

در نهایت روابط فوق را برای بیان نسبت‌های نمونه‌ای بتا-دوجمله‌ای $\frac{Z_i}{n_i}$ در شرایط تغییرنگار میدان تصادفی پنهان بتا، γ_Z ، به‌صورت

$$\begin{aligned} Var \left[\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right] &= E \left[\left(\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right)^2 \right] - E \left[\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right]^2 \\ &= \pi \left(\frac{n_i + n_j}{n_i n_j} \right) - \frac{n_i + n_j}{n_i n_j} (\sigma_Y^2 + \pi^2) + 2 \gamma_Y (s_i - s_j) \end{aligned}$$

می‌توان بازنویسی کرد. طبق تعریف، تغییرنگار به‌صورت $\frac{1}{2} Var \left[\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right]$ می‌باشد. بنابراین با ضرب معادله فوق در $\frac{1}{2}$ ، صورت نهایی تغییرنگار به‌صورت زیر حاصل می‌شود:

$$\frac{1}{2} Var \left[\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right] = \frac{1}{2} \left(\frac{n_i + n_j}{n_i n_j} \right) (\pi (1 - \pi) - \sigma_Y^2) + \gamma_Y (s_i - s_j).$$

به‌عبارت دیگر

$$\gamma_Z (s_i - s_j) = \frac{1}{2} \left(\frac{n_i + n_j}{n_i n_j} \right) (\pi (1 - \pi) - \sigma_Y^2) + \gamma_Y (s_i - s_j).$$

بنابراین تنها اختلاف تغییرنگار نظری نسبت‌های دوجمله‌ای مشاهده‌شده و توزیع بتا پنهان، توسط تابع ثابت پارامترهای بتا پنهان و اندازه نمونه است. پس

$$\gamma_Y(s_i - s_j) = \gamma_Z(s_i - s_j) - \frac{1}{4} \left(\frac{n_i + n_j}{n_i n_j} \right) (\pi(1 - \pi) - \sigma_Y^2). \quad (2.2)$$

علاوه بر این، امید ریاضی شرطی اختلاف توان دوم واریانس نیز برابر با

$$\begin{aligned} E \left[\text{Var} \left(\left[\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right]^2 \middle| Y \right) \right] &= E \left[\frac{1}{n_i} Y_i (1 - Y_i) \right] + E \left[\frac{1}{n_j} Y_j (1 - Y_j) \right] \\ &= \frac{1}{n_i} \pi - \frac{1}{n_j} (\sigma_Y^2 + \pi^2) + \frac{1}{n_j} \pi - \frac{1}{n_j} (\sigma_Y^2 + \pi^2) \\ &= \frac{n_i + n_j}{n_i n_j} (\pi(1 - \pi) - \sigma_Y^2) \end{aligned}$$

است.

۳.۲ برآوردگر تغییرنگار بتا-دوجمله‌ای

اکنون m مشاهده Z_i ، $i = 1, \dots, m$ ، از توزیع دوجمله‌ای با اندازه نمونه n_i در مکان s_i را در نظر بگیرید. یک برآورد تعدیل‌یافته برای تغییرنگار با استفاده از رابطه (۲.۲) به صورت زیر به دست می‌آید:

$$\begin{aligned} \gamma_Y(h) &= \frac{1}{N(h)} \sum_{i=1}^m \sum_{j=1}^m \frac{n_i n_j}{n_i + n_j} \left[\frac{1}{4} \left(\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right)^2 - \right. \\ &\quad \left. \frac{1}{4} \left(\frac{n_i + n_j}{n_i n_j} \right) (\hat{\pi}(1 - \hat{\pi}) - \hat{\sigma}_Y^2) \right] I_d(i, j) \sim h \end{aligned} \quad (3.2)$$

که در آن منظور از $I_d(i, j) \sim h$ تابع نشانگر برای نقاط i و j ای است که در فاصله h از هم قرار دارند. عدد $N(h)$ نیز تعداد نقاط همسایه در فاصله h است که بر حسب وزن‌های $\frac{n_i n_j}{n_i + n_j}$ برای همگن کردن تمام مکان‌های فضایی به صورت زیر تعریف می‌شود:

$$N(h) = \sum_{i,j} \frac{n_i n_j}{n_i + n_j} I_d(i, j) \sim h.$$

به عبارت دیگر این وزن تغییر تعداد مشاهدات در هر مکان فضایی را نشان می‌دهد. اگر همه مکان‌های فضایی اندازه نمونه یکسانی را داشته باشند، یعنی برای $i \neq j$ ، $n_i = n_j$ ، آن گاه به تمام نقاط مکانی وزن ثابت n داده می‌شود. مولفه تصحیح اریبی، یعنی

$$\left(\frac{n_i + n_j}{n_i n_j} \right) (\hat{\pi}(1 - \hat{\pi}) - \hat{\sigma}_Y^2)$$

به طور مستقیم از رابطه (۲.۲) تغییرنگار به دست می‌آید. همچنین $\hat{\pi}$ و $\hat{\sigma}_Y^2$ به ترتیب برآوردهای میانگین و واریانس توزیع بتای پنهان را نمایش می‌دهند.

در بسیاری از برنامه‌های کاربردی دانستن شکل توزیع بتا زیر بنایی (مقارن، چوله به چپ یا چوله به راست) مد نظر می‌باشد. بنابراین با جای گذاری پارمترهای شکل α و β در عبارت $(\hat{\pi}(1 - \hat{\pi}) - \hat{\sigma}_Y^2)$ ، خواهیم داشت

$$\begin{aligned}\pi(1 - \pi) - \sigma_Y^2 &= \pi - \pi^2 - \sigma_Y^2 \\ &= \frac{\alpha}{\alpha + \beta} - \frac{\alpha^2}{(\alpha + \beta)^2} - \frac{\alpha\beta}{(\alpha + \beta)^2(\alpha + \beta + 1)} \\ &= \frac{\alpha(\alpha + \beta)(\alpha + \beta + 1) - \alpha^2(\alpha + \beta + 1) - \alpha\beta}{(\alpha + \beta)^2(\alpha + \beta + 1)} \\ &= \frac{\alpha^2(\alpha + \beta + 1) + \alpha\beta(\alpha + \beta + 1) - \alpha\beta - \alpha^2(\alpha + \beta + 1)}{(\alpha + \beta)^2(\alpha + \beta + 1)} \\ &= \frac{\alpha\beta(\alpha + \beta)}{(\alpha + \beta)^2(\alpha + \beta + 1)} \\ &= \frac{\alpha\beta}{(\alpha + \beta)(\alpha + \beta + 1)}.\end{aligned}$$

بنابراین عبارت نهایی برآوردگر تغییرنگار با استفاده از پارمترهای شکل توزیع بتای پنهان به صورت زیر تبدیل می‌شود:

$$\begin{aligned}\gamma_Y(h) &= \frac{1}{2N(h)} \sum_{i=1}^m \sum_{j=1}^m \left[\frac{n_i n_j}{n_i + n_j} \left(\frac{Z_i}{n_i} - \frac{Z_j}{n_j} \right)^2 - \right. \\ &\quad \left. \left(\frac{\hat{\alpha}\hat{\beta}}{(\hat{\alpha} + \hat{\beta})(\hat{\alpha} + \hat{\beta} + 1)} \right) \right] I_d(i, j) \sim h.\end{aligned}\quad (4.2)$$

۴.۲ معادلات کریگینگ بتا-دوجمله‌ای

کریگینگ فضایی یک روش درونیایی است که در آن مقادیر پیش‌گویی ترکیب‌های خطی از مقادیر مشاهده‌شده، با وزن‌های تعریف‌شده به صورت تابعی از فاصله بین مکان مشاهده‌شده و مکان پیش‌گویی‌شده هستند. برای هر مکان فضایی جدید s_0 مقدار پیش‌گویی فرآیند بتای پنهان، $Y_0 = Y(s_0)$ ، می‌تواند بر حسب یک ترکیب خطی از نسبت‌های نمونه‌ای دوجمله‌ای $\frac{Z_i}{n_i}$ که در مکان‌های s_i مشاهده شده‌اند با وزن‌های λ_i ، $i = 1, \dots, m$ ، پیش‌گویی شود. یعنی

$$Y_0^* = \sum_{i=1}^m \lambda_i \frac{Z_i}{n_i}.$$

وزن‌های λ_i باید به گونه‌ای انتخاب شوند که پیش‌گوی Y_0 یک پیش‌گوی ناریب با کم‌ترین واریانس باشد. برای اطمینان از ناریبی پیش‌گو باید قید $\sum_{i=1}^m \lambda_i = 1$ را روی ضرایب اعمال کنیم. اگر مقدار پیش‌گویی Y_0^* ناریب باشد، آن‌گاه $E(Y_0^*) = \pi$ که میانگین توزیع بتای

پنهان است. بنابراین امید ریاضی شرطی در مکان جدید به صورت زیر است:

$$\begin{aligned} E[Y_*^*|Y] &= E\left[\sum_{i=1}^m \lambda_i \frac{Z_i}{n_i} | Y\right] \\ &= \sum_{i=1}^m \frac{\lambda_i}{n_i} E[Z_i | Y] \\ &= \sum_{i=1}^m \frac{\lambda_i}{n_i} [n_i Y_i] = \sum_{i=1}^m \lambda_i Y_i. \end{aligned}$$

لذا مقدار امید ریاضی Y_*^* برابر است با

$$\begin{aligned} E[Y_*^*] &= E\left[\sum_{i=1}^m \lambda_i Y_i\right] \\ &= \sum_{i=1}^m \lambda_i E[Y_i] \\ &= \pi \sum_{i=1}^m \lambda_i. \end{aligned}$$

در نتیجه، نااریبی کریگینگ مستلزم آن خواهد بود که $\sum_{i=1}^m \lambda_i = 1$ باشد. بنابراین اولین شرط کریگینگ بتا-دوجمله‌ای آن است که $\sum_{i=1}^m \lambda_i = 1$ باشد. برای بررسی شرط دوم کریگینگ، ابتدا امید ریاضی شرطی توان دوم خطای پیش‌گو را به صورت زیر به دست می‌آوریم:

$$\begin{aligned} E\left[(Y_*^* - Y_*)^2 | Y\right] &= E\left[\left(\sum_{i=1}^m \lambda_i \frac{Z_i}{n_i} - Y_*\right)^2 | Y\right] \\ &= E\left[\left\{\left(\sum_{i=1}^m \lambda_i \frac{Z_i}{n_i}\right)\left(\sum_{j=1}^m \lambda_j \frac{Z_j}{n_j}\right) - 2Y_* \sum_{i=1}^m \lambda_i \frac{Z_i}{n_i} + Y_*^2\right\} | Y\right] \\ &= \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j \frac{1}{n_i n_j} E[Z_i Z_j | Y] - 2Y_* \sum_{i=1}^m \frac{\lambda_i}{n_i} E[Z_i | Y] + Y_*^2 \\ &= \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j \frac{1}{n_i n_j} [\delta_{ij} n_i Y_i (1 - Y_i) + n_i n_j Y_i Y_j] - 2Y_* \sum_{i=1}^m \lambda_i Y_i + Y_*^2. \end{aligned}$$

امید ریاضی شرطی توان دوم خطای پیش‌گو را می‌توان به صورت زیر ساده کرد:

$$E\left[(Y_*^* - Y_*)^2 | Y\right] = \sum_{i=1}^m \frac{\lambda_i^2}{n_i} Y_i (1 - Y_i) + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j Y_i Y_j - 2Y_* \sum_{i=1}^m \lambda_i Y_i + Y_*^2.$$

بنابراین

$$\begin{aligned}
 E[(Y_{\circ}^* - Y_{\circ})] &= E \left[\sum_{i=1}^m \frac{\lambda_i^{\gamma}}{n_i} Y_i (1 - Y_i) + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j Y_i Y_j - \gamma Y_{\circ} \sum_{i=1}^m \lambda_i Y_i + Y_{\circ}^{\gamma} \right] \\
 &= \sum_{i=1}^m \frac{\lambda_i^{\gamma}}{n_i} E[Y_i (1 - Y_i)] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j E[Y_i Y_j] - \gamma \sum_{i=1}^m \lambda_i E[Y_{\circ} Y_i] + E[Y_{\circ}^{\gamma}] \\
 &= \sum_{i=1}^m \frac{\lambda_i^{\gamma}}{n_i} [E(Y_i) - E(Y_i^{\gamma})] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j [Cov(Y_i, Y_j) + E(Y_i) E(Y_j)] \\
 &\quad - \gamma \sum_{i=1}^m \lambda_i [Cov(Y_{\circ}, Y_i) + E(Y_{\circ}) E(Y_i)] + [Var(Y_{\circ}) + E(Y_{\circ}^{\gamma})] \\
 &= \sum_{i=1}^m \frac{\lambda_i^{\gamma}}{n_i} [\pi (1 - \pi) - \sigma_Y^{\gamma}] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j [C_Y(|s_i - s_j|) + \pi^{\gamma}] \\
 &\quad - \gamma \sum_{i=1}^m \lambda_i [C_Y(|s_{\circ} - s_i|) + \pi^{\gamma}] + \sigma_Y^{\gamma} + \pi^{\gamma} \\
 &= \sum_{i=1}^m \frac{\lambda_i^{\gamma}}{n_i} [\pi (1 - \pi) - \sigma_Y^{\gamma}] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j C_Y(|s_i - s_j|) + \pi^{\gamma} \\
 &\quad - \gamma \sum_{i=1}^m \lambda_i C_Y(s_{\circ} - s_i) - \gamma \pi^{\gamma} + \sigma_Y^{\gamma} + \pi^{\gamma} \\
 &= \sum_{i=1}^m \frac{\lambda_i^{\gamma}}{n_i} [\pi (1 - \pi) - \sigma_Y^{\gamma}] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j C_{ij} - \gamma \sum_{i=1}^m \lambda_i C_{\circ i} + \sigma_Y^{\gamma}
 \end{aligned}$$

که در آن

$$C_{ij} = C_Y(|s_i - s_j|).$$

در نتیجه واریانس پیش‌گو به‌صورت

$$\begin{aligned}
 Var(Y_{\circ}^* - Y_{\circ}) &= E[(Y_{\circ}^* - Y_{\circ})^{\gamma}] - E[Y_{\circ}^* - Y_{\circ}]^{\gamma} \\
 &= \left[\sum_{i=1}^m \frac{\lambda_i^{\gamma}}{n_i} [\pi (1 - \pi) - \sigma_Y^{\gamma}] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j C_{ij} - \gamma \sum_{i=1}^m C_{\circ i} + \sigma_Y^{\gamma} \right] - \circ
 \end{aligned}$$

به‌دست می‌آید. می‌دانیم که $Y_{\circ}^*$ نااریب است، یعنی $E[Y_{\circ}^* - Y_{\circ}]^{\gamma} = \circ$ ، بنابراین واریانس خطای پیش‌گو برابر است با

$$Var(Y_{\circ}^* - Y_{\circ}) = \sum_{i=1}^m \frac{\lambda_i^{\gamma}}{n_i} [\pi (1 - \pi) - \sigma_Y^{\gamma}] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j C_{ij} - \gamma \sum_{i=1}^m \lambda_i C_{\circ i} + \sigma_Y^{\gamma}.$$

اکنون با استفاده از روش لاگرانژ، واریانس خطای پیش‌گو را می‌نیمیم می‌کنیم. با مشتق‌گیری نسبت به λ_i داریم

$$\frac{\partial}{\partial \lambda_i} Var(Y_o^* - Y_o) = \frac{\partial}{\partial \lambda_i} \left\{ \sum_{i=1}^m \frac{\lambda_i^2}{n_i} [\pi(1-\pi) - \sigma_Y^2] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j C_{ij} - \right. \\ \left. 2 \sum_{i=1}^m \lambda_i C_{oi} + \sigma_Y^2 + 2\mu \left[\sum_{i=1}^m \lambda_i - 1 \right] \right\}$$

با مساوی قرار دادن مشتق با صفر دستگاه معادلات

$$\begin{cases} \circ &= \frac{2\lambda_i}{n_i} [\pi(1-\pi) - \sigma_Y^2] + 2 \sum_{j=1}^m \lambda_j C_{ij} - 2C_{oi} + 2\mu \\ \circ &= \frac{\lambda_i}{n_i} [\pi(1-\pi) - \sigma_Y^2] + \sum_{j=1}^m \lambda_j C_{ij} - C_{oi} + \mu \end{cases}$$

حل می‌کنیم. بنابراین معادله کریگینگ بتا-دوجمله‌ای به صورت زیر نتیجه می‌شود:

$$\frac{\lambda_i}{n_i} [\pi(1-\pi) - \sigma_Y^2] + \sum_{j=1}^m \lambda_j C_{ij} + \mu = C_{oi} \quad i \in 1, 2, \dots, n \\ \sum_{i=1}^m \lambda_i = 1.$$

معادلات کریگینگ را به راحتی می‌توان به صورت ماتریسی نوشت. در معادله بالا داریم

$$C\lambda = \underline{C}_o.$$

که در آن C ماتریس ضرایب، $\lambda = [\lambda_1, \dots, \lambda_n, \mu]'$ بردار وزن‌های کریگینگ و μ نیز ضریب لاگرانژ است. همچنین

$$\underline{C}_o = [C_{o1}, \dots, C_{on}, 1]'$$

بردار معادلات کریگینگ است.

معادله کریگینگ معمولی به صورت

$$\sum_{i=1}^n \lambda_i C_{ij} + \mu = C_{oi} \quad i \in 1, 2, \dots, n \\ \sum_{i=1}^m \lambda_i = 1$$

را در نظر بگیرید. ماتریس ضرایب کریگینگ بتا-دوجمله‌ای در مقایسه با کریگینگ معمولی پیچیده‌تر است، که این پیچیدگی اضافی از عبارت $\frac{\lambda_i}{n_i} [\pi(1-\pi) - \sigma_Y^2]$ ناشی شده است. در واقع در کریگینگ بتا-دوجمله‌ای، این جمله به عناصر قطر اصلی ماتریس ضرایب اضافه شده

است. بنابراین

$$C = \begin{bmatrix} \hat{C}_{11} + \left[\frac{\pi(1-\pi) - \sigma_Y^2}{n_1} \right] & \hat{C}_{12} & \dots & \hat{C}_{1n} & 1 \\ \hat{C}_{21} & +\hat{C}_{22} + \left[\frac{\pi(1-\pi) - \sigma_Y^2}{n_2} \right] & \dots & \cdot & 1 \\ \vdots & \vdots & \ddots & & \\ \hat{C}_{n1} & \vdots & \dots & \hat{C}_{nm} + \left[\frac{\pi(1-\pi) - \sigma_Y^2}{n_m} \right] & 1 \\ 1 & 1 & \dots & 1 & 0 \end{bmatrix}$$

به‌طوری‌که $\hat{C}_{ij} = \hat{\gamma}_Y(|s_i - s_j|)$ نشان‌دهنده تابع کوواریانس برآوردشده برای نقاط i و j می‌باشد. وزن‌های کریگینگ بتا-دوجمله‌ای با حل معادله ماتریسی $C\lambda = \underline{C}_0$ نتیجه می‌شوند. بنابراین

$$\lambda = C^{-1}\underline{C}_0.$$

برای انجام دادن کریگینگ بتا-دوجمله‌ای باید مجموعه معادلات برای هر مکان پیش‌گویی حل شوند، بنابراین زمانی که پیش‌گویی بیش از یک شبکه فضایی متراکم مد نظر باشد، ممکن است که محاسبات سنگین و زمان‌بر باشند.

همچنین واریانس خطای پیش‌گو را می‌توان در شرایط واریانس کریگینگ معمولی یعنی

$$Var_{OK} = \sigma_Y^2 - \sum_{i=1}^m \lambda_i C_{0i} - \mu$$

بازنویسی کرد. بنابراین با بازنویسی و ساده کردن عبارات داریم

$$\begin{aligned} Var(Y_o^* - Y_o) &= \sigma_Y^2 - \sum_{i=1}^m \lambda_i C_{0i} + \sum_{i=1}^m \frac{\lambda_i^2}{n_i} \left[\pi(1-\pi) - \sigma_Y^2 \right] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j C_{ij} \\ &= \sigma_Y^2 - \sum_{i=1}^m \lambda_i C_{0i} + \sum_{i=1}^m \frac{\lambda_i^2}{n_i} \left[\pi(1-\pi) - \sigma_Y^2 \right] + \sum_{i=1}^m \sum_{j=1}^m \lambda_i \lambda_j C_{ij} \\ &\quad - \sum_{i=1}^m \lambda_i C_{0i} \\ &= \sigma_Y^2 - \sum_{i=1}^m \lambda_i C_{0i} + \sum_{i=1}^m \lambda_i \left\{ \frac{\lambda_i}{n_i} \left[\pi(1-\pi) - \sigma_Y^2 \right] + \sum_{j=1}^m \lambda_j C_{ij} - C_{0i} \right\} \\ &= \sigma_Y^2 - \sum_{i=1}^m \lambda_i C_{0i} - \sum_{i=1}^m \lambda_i \mu \\ &= \sigma_Y^2 - \sum_{i=1}^m \lambda_i C_{0i} - \mu \end{aligned}$$

از رابطه بالا می‌توان چنین استنباط کرد که پیچیدگی اضافی نسبت‌های نمونه‌ای توزیع بتا-دوجمله‌ای، از واریانس خطای پیش‌گو ناشی نشده و این پیچیدگی فقط در ماتریس کوواریانس وجود دارد.

در فصل بعد عملکرد برآوردگر تغییرنگار و کریگینگ بتا-دوجمله‌ای را مورد ارزیابی قرار می‌دهیم.

فصل ۳

دقت برآوردیابی و پیش‌گویی مدل بتا-دوجمله‌ای فضایی

۱.۳ مقدمه

در این فصل ابتدا به تشریح نحوه تولید داده‌های بتا-دوجمله‌ای همبسته فضایی با استفاده از روش تبدیل معکوس می‌پردازیم. سپس سه روش مختلف را برای برآورد پارامترهای تغییرنگار معرفی کرده، که پس از بررسی دقت آن‌ها بهترین روش برای برآورد پارامترهای کریگینگ بتا-دوجمله‌ای را انتخاب می‌کنیم. علاوه بر این، دقت مدل در پیش‌گویی را با معیارهای مختلف مورد ارزیابی قرار می‌دهیم.

۱.۱.۳ روش تبدیل معکوس برای شبیه‌سازی داده‌ها

یاه و شیمولی (۲۰۰۷)، روش تبدیل معکوس را به‌عنوان روشی کارا برای تولید متغیر تصادفی پواسون چند متغیره معرفی کردند. در بعضی از متون این روش را تحت عنوان تبدیل NORTA یعنی "نرمال به هر چیزی" می‌نامند. اساس این روش تبدیل از توزیع نرمال چندمتغیره به هر توزیع احتمالی دیگر با استفاده از تابع توزیع تجمعی است.

فرآیند شبیه‌سازی روش تبدیل معکوس شامل سه مرحله به‌صورت زیر است:

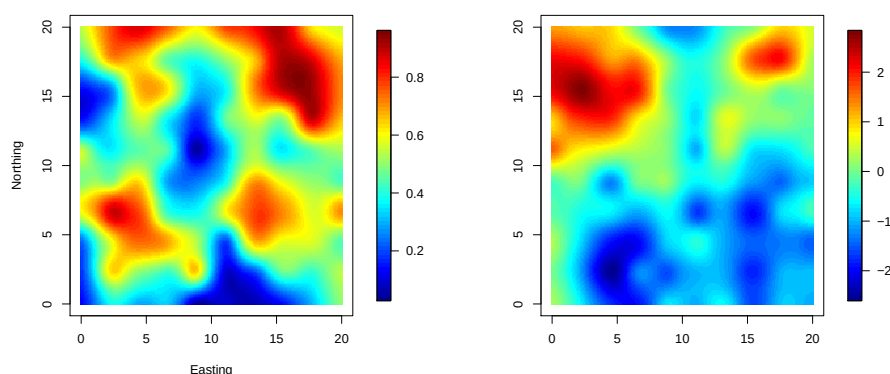
۱. ابتدا بردار توزیع نرمال p بعدی X^N ، با میانگین $\mu = 0$ ، و ماتریس همبستگی R^N را تولید می‌کنیم.

۲. تابع توزیع تجمعی را برای هر مقدار بردار توزیع نرمال X_i^N ، $i \in 1, 2, \dots, p$ ، به صورت $\Phi(X_i)$ محاسبه می‌کنیم.

۳. در نهایت معکوس تابع توزیع تجمعی (تابع چندک) توزیع دلخواه، مثلاً پواسون با پارمتر λ_i ، را به صورت

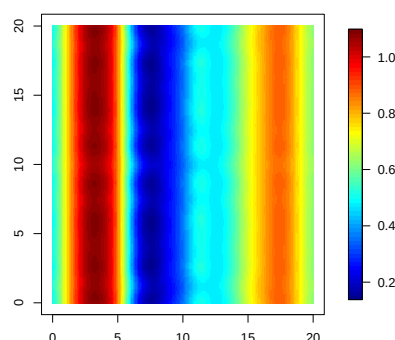
$$X_i^D = \Xi^{-1}(\Phi(X_i^N))$$

به دست می‌آوریم. در نتیجه X^D یک بردار تصادفی p بعدی از توزیع دلخواه با ماتریس همبستگی R^D است.



(ب) میدان تصادفی بتا

(آ) میدان تصادفی نرمال



(ج) میدان نسبت‌های دوجمله‌ای

شکل ۱۰۳: میدان‌های شبیه‌سازی شده با روش تبدیل معکوس

روش تبدیل معکوس در مقایسه با سایر روش‌ها، ساده و انعطاف‌پذیر است. در این روش

هیچ گونه محدودیتی بر روی ماتریس همبستگی داده‌های نرمال R^D وجود ندارد. همچنین یاها و شیمولی پیشنهاد کردند از متغیرهای تصادفی نرمال استاندارد استفاده شود که علت این انتخاب فقط سادگی آن می‌باشد. مزیت دیگر این روش آن است که بر روی اکثر توزیع‌ها به خوبی عمل می‌کند. البته این روش وقتی که توزیع مورد نظر پیوسته باشد، عملکرد بهتری دارد. همچنین یاها و شیمولی اشاره کردند که سختی‌هایی در تعریف معکوس تابع توزیع تجمعی توزیع‌های گسسته مثل پواسون وجود دارند و راه حل‌هایی را برای آن‌ها پیشنهاد دادند. برای اطلاعات بیشتر تر به یاها و شیمولی (۲۰۰۷) مراجعه کنید.

۲.۱.۳ شبیه‌سازی میدان فضایی بتا-دوجمله‌ای

همان‌طور که در بخش قبل عنوان شد، از روش تبدیل معکوس برای تولید نمونه از توزیع‌های مختلف می‌توان استفاده کرد. بنابراین از این روش برای تولید داده‌های بتا-دوجمله‌ای همبسته فضایی استفاده می‌کنیم که مراحل آن به صورت زیر هستند:

۱. ابتدا میدان تصادفی نرمال $X(s)$ با ماتریس کوواریانس فضایی Σ_X و تابع کوواریانس فضایی γ_X را شبیه‌سازی می‌کنیم. در این مطالعه شبیه‌سازی، نقاط مکانی، گره‌های یک شبکه فضایی 20×20 در نظر گرفته شدند. برای تابع کوواریانس فضایی نیز یک مدل همسان‌گرد کروی با پارامترهای واقعی اثر قطعه‌ای $\tau_X^2 = 0$ ، دامنه فضایی $\phi_X = 5, 10, 15$ و ازاره فضایی $\sigma_X^2 = 1$ در نظر گرفته شد. بنابراین

$$X(s) \sim N_m(\mu_X = 0, \Sigma_X)$$

$$\gamma_X(h) = \begin{cases} \tau_X^2 + \sigma_X^2 \left(\frac{\tau_X^2 h}{\tau_X^2 \phi_X} - \frac{h^2}{\tau_X^2 \phi_X^2} \right) & 0 < h < \phi_X \\ \tau_X^2 + \sigma_X^2 & h \geq \phi_X \end{cases}$$

۲. برای هر نقطه مکانی در شبکه فضایی، تابع توزیع تجمعی نرمال استاندارد، یعنی $\Phi_X(X_i)$ ، را برای تبدیل میدان تصادفی نرمال به میدان تصادفی بتا محاسبه می‌کنیم. اگر تابع چندک (معکوس تابع توزیع تجمعی) توزیع بتا $Y_i \sim \text{Beta}(\alpha, \beta)$ را با $\Xi_Y^{-1}(\cdot)$ نشان دهیم، آن‌گاه

$$\Xi_Y^{-1}(\Phi(X_i)) = Y_i.$$

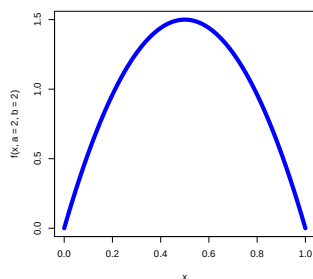
چون هر دو توزیع پیوسته هستند، تبدیل انجام‌شده یک به یک خواهد بود. توزیع بتا دارای ساختار وابستگی Σ_Y است، که این ساختار هم از یک مدل کروی با اثر قطعه‌ای $\tau_Y^2 = 0$ پیروی می‌کند. پارامتر دامنه نیز مشابه مدل میدان تصادفی نرمال است. فقط پارامتر ازاره فضایی برای میدان تصادفی بتا برابر با واریانس توزیع بتا به صورت زیر خواهد بود.

$$\sigma_Y^2 = \frac{\alpha\beta}{(\alpha + \beta)^2 (\alpha + \beta + 1)}$$

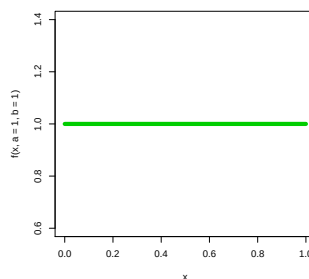
۳. یک زیرمجموعه m تایی از مکان‌های فضایی با اندازه نمونه تصادفی n_i را به صورت تصادفی انتخاب می‌کنیم. سپس در هر مکان نمونه‌گیری، نمونه تصادفی دوجمله‌ای $Z_i|Y_i \sim \text{bin}(n_i, Y_i)$ تولید می‌شود. بنابراین متغیر تصادفی Z_i دارای توزیع بتا-دوجمله‌ای همبسته فضایی است. شکل ۱.۳ فرآیند تولید داده‌های بتا-دوجمله‌ای همبسته فضایی را نمایش می‌دهد.

۳.۱.۳ مشخصات داده‌های شبیه‌سازی شده بتا-دوجمله‌ای فضایی

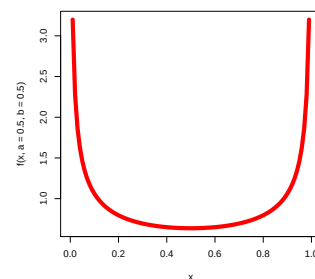
برای ارزیابی عملکرد مدل بتا-دوجمله‌ای پیشنهادی در برآورد پارامترها و پیش‌گویی فضایی، از یک مطالعه شبیه‌سازی استفاده کردیم. به منظور نمایش انعطاف توزیع‌های بتا که شامل توزیع‌های متقارن و چوله می‌باشد، در این پایان‌نامه شش توزیع بتا مختلف را انتخاب کردیم. توابع چگالی احتمال توزیع‌های انتخاب شده در شکل ۲.۳ نشان داده شده‌اند. همچنین جدول ۱.۳ میانگین و واریانس این شش توزیع بتا را نشان می‌دهد.



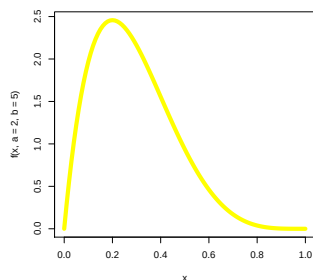
Beta(۲،۲) (ج)



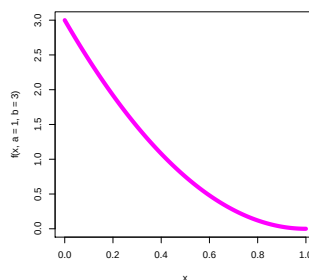
Beta(۱،۱) (ب)



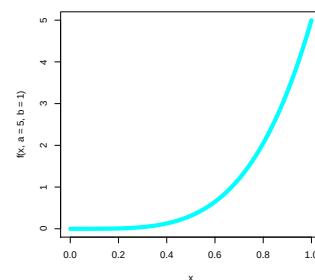
Beta(۰.۵،۰.۵) (آ)



Beta(۲،۵) (و)



Beta(۱،۳) (ه)



Beta(۵،۱) (د)

شکل ۲.۳: توابع چگالی بتا

جدول ۱.۳: میانگین و واریانس توزیع‌های بتا

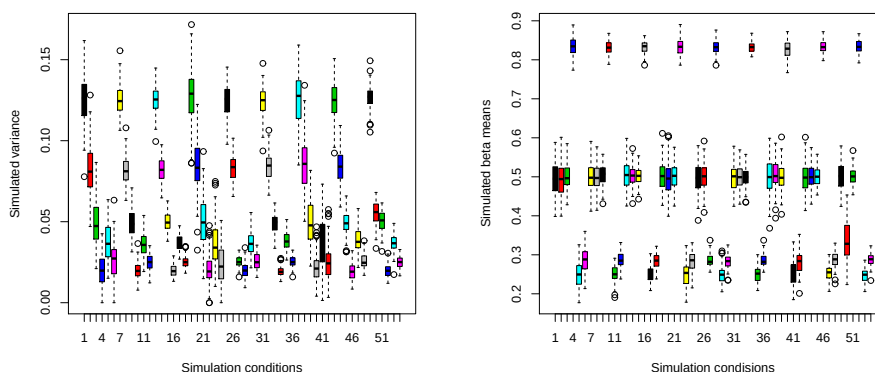
موارد	پارامترهای بتا	میانگین	واریانس
۱	$Beta(0/5, 0/5)$	۰/۵	۰/۱۲۵
۲	$Beta(1, 1)$	۰/۵	۰/۰۸۳
۳	$Beta(2, 2)$	۰/۵	۰/۰۵۰
۴	$Beta(5, 1)$	۰/۸۳۳	۰/۰۲۰
۵	$Beta(1, 3)$	۰/۲۵	۰/۰۳۸
۶	$Beta(2, 5)$	۰/۲۸۶	۰/۰۲۶

برای مطالعه شبیه‌سازی، ۵۴ سناریوی مختلف تولید داده را تعریف کردیم؛ برای هر توزیع بتا سه اندازه نمونه در نظر گرفتیم: محدوده بین ۲:۵ برای اندازه نمونه کوچک، محدوده بین ۱۰:۲۰ برای اندازه نمونه متوسط و محدوده بین ۵۰:۱۰۰ برای اندازه نمونه بزرگ تعریف کردیم. پارامتر دامنه فضایی را نیز سه مقدار ۵، ۱۰ و ۱۵ در نظر گرفتیم. مجموع ۵۴ سناریو شبیه‌سازی شده در جدول ۲.۳ نشان داده شده‌اند. تمامی داده‌ها را بدون اثر قطعه‌ای، به صورت هموار، بر روی شبکه منظم 20×20 تولید کردیم. بنابراین داده‌ها باید توزیع بتا واقعی را نمایش دهند و هرگونه اثر قطعه‌ای که در نسبت‌های نمونه‌ای دوجمله‌ای مشاهده شود فقط از نمونه‌گیری دوجمله‌ای ناشی شده است. همچنین برای هر سناریو ۱۰۰ مجموعه داده شبیه‌سازی کردیم.

جدول ۲.۳: سناریوهای مختلف مطالعه شبیه‌سازی

اندازه نمونه	دامنه	توزیع بتا	شماره شبیه‌سازی	اندازه نمونه	دامنه	توزیع بتا	شماره شبیه‌سازی
Medium(۱۰:۲۰)	۱۰	Beta(۵, ۱)	۲۸	small(۲:۵)	۵	Beta(۰/۵, ۰/۵)	۱
Medium(۱۰:۲۰)	۱۰	Beta(۱, ۳)	۲۹	Small(۲:۵)	۵	Beta(۱, ۱)	۲
Medium(۱۰:۲۰)	۱۰	Beta(۲, ۵)	۳۰	Small(۲:۵)	۵	Beta(۲, ۲)	۳
Large(۵۰:۱۰۰)	۱۰	Beta(۰/۵, ۰/۵)	۳۱	Small(۲:۵)	۵	Beta(۵, ۱)	۴
Large(۵۰:۱۰۰)	۱۰	Beta(۱, ۱)	۳۲	Small(۲:۵)	۵	Beta(۱, ۳)	۵
Large(۵۰:۱۰۰)	۱۰	Beta(۲, ۲)	۳۳	Small(۲:۵)	۵	Beta(۲, ۵)	۶
Large(۵۰:۱۰۰)	۱۰	Beta(۵, ۱)	۳۴	Medium(۱۰:۲۰)	۵	Beta(۰/۵, ۰/۵)	۷
Large(۵۰:۱۰۰)	۱۰	Beta(۱, ۳)	۳۵	Medium(۱۰:۲۰)	۵	Beta(۱, ۱)	۸
Large(۵۰:۱۰۰)	۱۰	Beta(۲, ۵)	۳۶	Medium(۱۰:۲۰)	۵	Beta(۲, ۲)	۹
Small(۲:۵)	۱۵	Beta(۰/۵, ۰/۵)	۳۷	Medium(۱۰:۲۰)	۵	Beta(۵, ۱)	۱۰
Small(۲:۵)	۱۵	Beta(۱, ۱)	۳۸	Medium(۱۰:۲۰)	۵	Beta(۱, ۳)	۱۱
Small(۲:۵)	۱۵	Beta(۲, ۲)	۳۹	Medium(۱۰:۲۰)	۵	Beta(۲, ۵)	۱۲
Small(۲:۵)	۱۵	Beta(۵, ۱)	۴۰	Large(۵۰:۱۰۰)	۵	Beta(۰/۵, ۰/۵)	۱۳
Small(۲:۵)	۱۵	Beta(۱, ۳)	۴۱	Large(۵۰:۱۰۰)	۵	Beta(۱, ۱)	۱۴
Small(۲:۵)	۱۵	Beta(۲, ۵)	۴۲	Large(۵۰:۱۰۰)	۵	Beta(۲, ۲)	۱۵
Medium(۰:۲۰)	۱۵	Beta(۰/۵, ۰/۵)	۴۳	Large(۵۰:۱۰۰)	۵	Beta(۵, ۱)	۱۶
Medium(۱۰:۲۰)	۱۵	Beta(۱, ۱)	۴۴	Large(۵۰:۱۰۰)	۵	Beta(۱, ۳)	۱۷
Medium(۱۰:۲۰)	۱۵	Beta(۲, ۲)	۴۵	Large(۵۰:۱۰۰)	۵	Beta(۲, ۵)	۱۸
Medium(۱۰:۲۰)	۱۵	Beta(۵, ۱)	۴۶	Small(۲:۵)	۱۰	Beta(۰/۵, ۰/۵)	۱۹
Medium(۱۰:۲۰)	۱۵	Beta(۱, ۳)	۴۷	Small(۲:۵)	۱۰	Beta(۱, ۱)	۲۰
Medium(۱۰:۲۰)	۱۵	Beta(۲, ۵)	۴۸	Small(۲:۵)	۱۰	Beta(۲, ۲)	۲۱
Large(۵۰:۱۰۰)	۱۵	Beta(۰/۵, ۰/۵)	۴۹	Small(۲:۵)	۱۰	Beta(۵, ۱)	۲۲
Large(۵۰:۱۰۰)	۱۵	Beta(۱, ۱)	۵۰	Small(۲:۵)	۱۰	Beta(۱, ۳)	۲۳
Large(۵۰:۱۰۰)	۱۵	Beta(۲, ۲)	۵۱	Small(۲:۵)	۱۰	Beta(۲, ۵)	۲۴
Large(۵۰:۱۰۰)	۱۵	Beta(۵, ۱)	۵۲	Medium(۱۰:۲۰)	۱۰	Beta(۰/۵, ۰/۵)	۲۵
Large(۵۰:۱۰۰)	۱۵	Beta(۱, ۳)	۵۳	Medium(۱۰:۲۰)	۱۰	Beta(۱, ۱)	۲۶
Large(۵۰:۱۰۰)	۱۵	Beta(۲, ۵)	۵۴	Medium(۱۰:۲۰)	۱۰	Beta(۲, ۲)	۲۷

نمودارهای جعبه‌ای شکل ۳.۳ میانگین و واریانس ۱۰۰ مجموعه داده شبیه‌سازی شده را تحت سناریوهای مختلف شبیه‌سازی نمایش می‌دهند. همان‌طور که در شکل مشاهده می‌شود هر یک از سناریوهای توزیع بتا به میانگین واقعی ($\circ/۲۸۶, \circ/۲۵, \circ/۸۳۳, \circ/۵, \circ/۵, \circ/۵$) نزدیک هستند. همچنین با افزایش دامنه فضایی در میانگین توزیع بتای شبیه‌سازی شده، تغییرات بیش‌تری ایجاد می‌شود. وابستگی فضایی شدید می‌تواند مقادیر میدان تصادفی را افزایش یا کاهش دهد. در واقع مکان‌های فضایی با وابستگی شدید، مکان‌های همسایه را بیش‌تر تحت تاثیر قرار می‌دهند، که این پدیده بیش‌تر برای توزیع‌های بتای متقارن تجلی پیدا می‌کند. با توجه به این که نمونه‌های دوجمله‌ای هنوز گرفته نشده‌اند بنابراین هیچ‌گونه اختلاف قابل تشخیصی بین اندازه‌های نمونه مشاهده نمی‌شود. به‌عنوان مثال، سناریو شبیه‌سازی ۷،۱ و ۱۳ در یک سطح از مدل موثر هستند. همچنین واریانس توزیع بتای شبیه‌سازی شده در این سه سناریو به واریانس واقعی توزیع بتا نزدیک هستند. مشابه میانگین، واریانس نیز به دامنه فضایی وابسته است، به‌طوری که با افزایش دامنه فضایی تغییرات واریانس نیز بیش‌تر می‌شود.



(آ) میانگین توزیع بتا شبیه‌سازی شده (ب) واریانس توزیع بتا شبیه‌سازی شده

شکل ۳.۳: میانگین و واریانس شبیه‌سازی شده

۴.۱.۳ برآورد پارامترهای توزیع بتا

برای مواردی که تعداد مشاهدات در تمامی مکان‌ها یکسان هستند، چندین مطالعه خوب برآورد برای پارامترها شامل $\hat{\pi}$ و $\hat{\sigma}_Y^2$ یا $(\hat{\alpha}, \hat{\beta})$ انجام شده‌اند. با توجه به این که پارامترهای شکل α و β به‌طور مستقیم در معادله (۴.۲) تغییرنگار بتا-دوجمله‌ای (فصل دوم) مورد استفاده قرار می‌گیرند، بنابراین برآورد دقیق پارامترهای شکل برای برازش مدل بتا-دوجمله‌ای ضروری است. با توجه به این که اندازه نمونه n_i در تمامی مکان‌ها یکسان نیست، بنابراین برآوردهایی که پیش از این پیشنهاد شده‌اند، مناسب نیستند. تربیتی (۱۹۹۴) برای اندازه نمونه نابرابر استفاده از روش‌های ماکسیمم درست‌نمایی را پیشنهاد داد، که برای ارزیابی مدل مورد استفاده

قرار خواهد گرفت. کریگینگ بتا-دوجمله‌ای یک مزیت منحصر به فرد را برای مدل‌سازی نمونه‌های کوچک فراهم می‌کند. برآورد ماکسیمم درست‌نمایی پارامترهای توزیع بتا، برای اندازه نمونه نابرابر مورد بررسی قرار نگرفته است. برای اطمینان از مناسب بودن این روش، یک مطالعه شبیه‌سازی را برای ارزیابی پارامترهای شکل برآوردشده انجام دادیم. نسبت نمونه‌ای $\frac{Z_i}{n_i}$ را به عنوان مقادیر مشاهده‌شده برای برآورد $\hat{\alpha}$ و $\hat{\beta}$ در نظر گرفتیم.

برای هر شش توزیع بتا، ۱۰۰۰ نمونه تصادفی بتا-دوجمله‌ای را در ۱۰۰ مکان نمونه‌گیری، روی شبکه 20×20 تولید کردیم. برآورد IWLS با استفاده از بسته VGAM در نرم‌افزار R برای برآورد پارامترهای توزیع بتا به دست آمدند. جدول‌های ۳.۳، ۴.۳ و ۵.۳ به ترتیب شبیه‌سازی‌ها را برای دامنه‌های فضایی ۵، ۱۰ و ۱۵ نشان می‌دهند. با توجه به نتایج جدول‌ها مشاهده می‌کنیم که پارامترهای برآوردشده بسیار نزدیک به پارامترهای شبیه‌سازی واقعی برای تمام توزیع‌های بتا هستند. البته برای نمونه‌های بزرگ اندکی کم‌برآوردی^۱ وجود دارد که خیلی قابل توجه نیست، و ممکن است علت آن وجود اریبی کوچک در میانگین و میانه برآورد پارامترها باشد.

جدول ۳.۳: برآورد پارامترهای شکل $\hat{\alpha}$ و $\hat{\beta}$ برای هر یک از توزیع‌های بتا با دامنه ۵

$\hat{\beta}_{Large}$	$\hat{\alpha}_{Large}$	$\hat{\beta}_{Medium}$	$\hat{\alpha}_{Medium}$	$\hat{\beta}_{Small}$	$\hat{\alpha}_{Small}$	توزیع بتا	
۰/۴۹۴	۰/۴۹۴	۰/۴۹۳	۰/۴۹۳	۰/۴۹۹	۰/۴۹۹	میانه	$Beta(0.5, 0.5)$
۰/۴۹۷	۰/۴۹۷	۰/۴۹۴	۰/۴۹۴	۰/۵۰۲	۰/۵۰۳	میانگین	
۰/۰۲۲	۰/۰۲۲	۰/۰۳۳	۰/۰۳۲	۰/۰۵۷	۰/۰۵۸	انحراف معیار	
۰/۹۹۵	۰/۹۹۵	۰/۹۹۴	۰/۹۹۳	۰/۹۷۸	۰/۹۸۵	میانه	$Beta(1, 1)$
۰/۹۹۷	۰/۹۹۷	۰/۹۹۴	۰/۹۹۵	۰/۹۹۸	۰/۹۹۸	میانگین	
۰/۰۲۸	۰/۰۲۸	۰/۰۵۷	۰/۰۵۸	۰/۱۲۶	۰/۱۲۴	انحراف معیار	
۱/۹۹۸	۱/۹۹۸	۱/۹۹۲	۱/۹۹۰	۱/۹۹۹	۲/۰۱۱	میانگین	$Beta(2, 2)$
۱/۹۹۹	۱/۹۹۹	۱/۹۹۹	۲/۰۰۱	۲/۰۳۶	۲/۰۳۸	میانه	
۰/۰۵۲	۰/۰۵۰	۰/۱۲۱	۰/۱۲۲	۰/۳۱۶	۰/۳۱۸	انحراف معیار	
۰/۹۹۶	۴/۹۸۹	۰/۹۹۸	۴/۹۷۳	۰/۹۹۸	۴/۹۵۳	میانه	$Beta(5, 1)$
۱/۰۰۰	۵/۰۰۵	۱/۰۰۵	۵/۰۲۳	۱/۰۳۹	۵/۱۹۷	میانگین	
۰/۰۴۹	۰/۲۶۵	۰/۱۰۰	۰/۵۱۳	۰/۲۶۶	۱/۳۵۱	انحراف معیار	
۲/۹۶۸	۰/۹۸۹	۲/۹۷۵	۰/۹۹۰	۳/۰۰۴	۱/۰۰۵	میانه	$Beta(1, 3)$
۲/۹۸۷	۰/۹۹۴	۳/۰۰۲	۱/۰۰۱	۳/۰۹۳	۱/۰۳۲	میانگین	
۰/۱۴۸	۰/۰۴۷	۰/۳۰۸	۰/۱۰۵	۰/۶۶۷	۰/۲۲۴	انحراف معیار	
۴/۹۹۳	۱/۹۹۶	۵/۰۱۹	۲/۰۰۸	۵/۰۳۹	۲/۰۰۷	میانگین	$Beta(2, 5)$
۵/۰۰۵	۲/۰۰۱	۵/۰۴۱	۲/۰۱۷	۵/۲۳۸	۲/۰۹۵	میانه	
۰/۱۷۰	۰/۰۶۷	۰/۴۴۴	۰/۱۷۶	۱/۳۱۷	۰/۵۲۱	انحراف معیار	

^۱ Underestimation

جدول ۴.۳: برآورد پارامترهای شکل $\hat{\alpha}$ و $\hat{\beta}$ برای هریک از توزیع‌های بتا با دامنه ۱۰

$\hat{\beta}_{Large}$	$\hat{\alpha}_{Large}$	$\hat{\beta}_{Medium}$	$\hat{\alpha}_{Medium}$	$\hat{\beta}_{Small}$	$\hat{\alpha}_{Small}$	توزیع بتا	
۰/۴۸۵	۰/۴۸۵	۰/۴۸۳	۰/۴۸۳	۰/۴۸۳	۰/۴۸۳	میانه	$Beta(۰/۵, ۰/۵)$
۰/۴۸۹	۰/۴۸۹	۰/۴۸۷	۰/۴۸۸	۰/۴۹۰	۰/۴۸۹	میانه	
۰/۰۲۸	۰/۰۲۸	۰/۰۴۲	۰/۰۴۱	۰/۰۷۰	۰/۰۷۰	انحراف معیار	
۰/۹۸۹	۰/۹۸۷	۰/۹۷۹	۰/۹۸۱	۰/۹۷۵	۰/۹۷۳	میانه	$Beta(1, 1)$
۰/۹۹۱	۰/۹۹۱	۰/۹۸۵	۰/۹۸۵	۰/۹۸۷	۰/۹۸۸	میانگین	
۰/۰۳۰	۰/۰۳۰	۰/۰۶۳	۰/۰۶۳	۰/۱۲۹	۰/۱۳۰	انحراف معیار	
۱/۹۸۷	۱/۹۹۱	۱/۹۷۷	۱/۹۷۷	۱/۹۶۷	۱/۹۶۸	میانه	$Beta(2, 2)$
۱/۹۹۱	۱/۹۹۱	۱/۹۸۴	۱/۹۸۳	۲/۰۰۲	۲/۰۰۳	میانه	
۰/۰۵۰	۰/۰۵۱	۰/۱۲۶	۰/۱۲۴	۰/۳۲۸	۰/۳۳۵	انحراف معیار	
۰/۹۸۸	۴/۹۳۴	۰/۹۹۰	۴/۹۵۸	۰/۹۹۱	۴/۹۵۹	میانه	$Beta(5, 1)$
۰/۹۹۵	۴/۹۷۶	۱/۰۰۵	۵/۰۲۳	۱/۰۵۳	۵/۲۷۷	میانگین	
۰/۰۵۸	۰/۳۱۱	۰/۱۲۶	۰/۶۴۳	۰/۲۹۰	۱/۴۹۶	انحراف معیار	
۲/۹۶۷	۰/۹۸۸	۲/۹۴۸	۰/۹۸۳	۳/۰۲۹	۱/۰۱۰	میانه	$Beta(1, 3)$
۲/۹۸۵	۰/۹۹۵	۲/۹۸۶	۰/۹۹۵	۳/۱۱۱	۱/۰۴۱	میانگین	
۰/۱۴۲	۰/۰۴۵	۰/۲۸۲	۰/۰۹۶	۰/۶۱۹	۰/۲۱۰	انحراف معیار	
۴/۹۸۰	۱/۹۹۲	۴/۹۸۸	۱/۹۹۱	۵/۰۵۰	۲/۰۳۰	میانه	$Beta(2, 5)$
۴/۹۸۴	۱/۹۹۳	۵/۰۰۸	۲/۰۰۴	۵/۲۵۶	۲/۱۰۱	میانگین	
۰/۱۸۹	۰/۰۷۳	۰/۴۷۱	۰/۱۹۱	۱/۳۱۲	۰/۵۱۸	انحراف معیار	

جدول ۵.۳: برآورد پارامترهای شکل $\hat{\alpha}$ و $\hat{\beta}$ برای هریک از توزیع‌های بتا با دامنه ۱۵

$\hat{\beta}_{Large}$	$\hat{\alpha}_{Large}$	$\hat{\beta}_{Medium}$	$\hat{\alpha}_{Medium}$	$\hat{\beta}_{Small}$	$\hat{\alpha}_{Small}$	توزیع بتا	
۰/۴۸۲	۰/۴۸۱	۰/۴۷۲	۰/۴۷۳	۰/۴۷۳	۰/۴۶۹	میانه	$Beta(۰/۵, ۰/۵)$
۰/۴۸۶	۰/۴۸۶	۰/۴۷۸	۰/۴۷۶	۰/۴۷۶	۰/۴۷۸	میانگین	
۰/۰۳۱	۰/۰۳۰	۰/۰۴۶	۰/۰۴۵	۰/۰۷۰	۰/۰۷۵	انحراف معیار	
۰/۹۸۲	۰/۹۸۲	۰/۹۷۲	۰/۹۷۱	۰/۹۵۵	۰/۹۵۶	میانه	$Beta(1, 1)$
۰/۹۸۷	۰/۹۸۷	۰/۹۷۹	۰/۹۷۶	۰/۹۷۰	۰/۹۷۱	میانگین	
۰/۰۳۳	۰/۰۳۳	۰/۰۶۸	۰/۰۶۹	۰/۱۳۶	۰/۱۳۵	انحراف معیار	
۱/۹۹۸	۱/۹۹۸	۱/۹۹۲	۱/۹۹۰	۱/۹۹۹	۲/۰۱۱	میانه	$Beta(2, 2)$
۱/۹۹۹	۱/۹۹۹	۱/۹۹۹	۲/۰۰۱	۲/۰۳۶	۲/۰۳۸	میانه	
۰/۰۵۲	۰/۰۵۰	۰/۱۲۱	۰/۱۲۲	۰/۳۱۶	۰/۳۱۸	انحراف معیار	
۰/۹۷۹	۴/۹۰۴	۰/۹۹۸	۴/۹۳۹	۱/۰۱۱	۵/۰۷۷	میانه	$Beta(5, 1)$
۰/۹۹۰	۴/۹۵۰	۱/۰۰۵	۵/۰۲۹	۱/۰۶۵	۵/۳۲۰	میانگین	
۰/۰۶۱	۰/۳۲۹	۰/۱۴۰	۰/۷۱۱	۰/۳۹۰	۲/۰۶۶	انحراف معیار	
۲/۹۵۴	۰/۹۸۲	۲/۹۵۲	۰/۹۸۴	۲/۹۵۸	۰/۹۸۹	میانه	$Beta(1, 3)$
۲/۹۷۴	۰/۹۹۱	۲/۹۹۹	۱/۰۰۰	۳/۰۵۹	۱/۰۲۳	میانه	
۰/۱۵۲	۰/۰۴۷	۰/۳۱۹	۰/۱۰۷	۰/۶۴۶	۰/۲۲۲	انحراف معیار	
۴/۹۷۲	۱/۹۸۸	۴/۹۳۵	۱/۹۷۲	۵/۰۳۴	۲/۰۱۲	میانه	$Beta(2, 5)$
۴/۹۸۱	۱/۹۹۲	۴/۹۶۳	۱/۹۸۷	۵/۲۶۸	۲/۱۰۶	میانگین	
۰/۱۹۱	۰/۰۷۳	۰/۴۷۵	۰/۱۹۲	۱/۳۹۸	۰/۵۶۱	انحراف معیار	

۵.۱.۳ ویژگی‌های تغییرنگار بتا-دوجمله‌ای

با توجه به این که متغیر پاسخ در این پایان‌نامه به صورت نسبت است، بنابراین هر یک از پارامترهای تغییرنگار شامل دامنه، ازاره و اثر قطعه‌ای را با توجه به نسبت‌های مشاهده شده $\frac{Z_i}{n_i}$ بررسی می‌کنیم. همان‌طور که مطرح کردیم توزیع دوجمله‌ای شمارشی دارای توزیع کناری بتا-دوجمله‌ای است، یعنی

$$Z_i \sim \text{beta-bin}(n_i, Y_i, \alpha, \beta).$$

پیش از ارزیابی عملکرد مدل کریگینگ بتا-دوجمله‌ای، برخی از ویژگی‌های روابط فضایی نسبت‌های نمونه‌ای مشاهده شده $\frac{Z_i}{n_i}$ باید اثبات شوند. با توجه به این که هر دو توزیع نرمال و بتا پیوسته هستند، بنابراین هیچ‌گونه تغییر اضافی در فرآیند تبدیل از میدان تصادفی نرمال به میدان تصادفی بتا به مدل اضافه نمی‌شود و هرگونه تغییر اضافی در اثر قطعه‌ای در سطح نسبت‌های دوجمله‌ای رخ داده است. برای اندازه نمونه ثابت در هر مکان فضایی، روابط بین واریانس میدان تصادفی بتای پنهان، $Var(Y)$ ، و واریانس میدان تصادفی بتا-دوجمله‌ای، $Var(Z)$ ، به صورت زیر می‌باشد:

$$Var(Y) = \frac{\alpha\beta}{(\alpha + \beta)^2 (\alpha + \beta + 1)}$$

$$Var(Z) = \frac{n\alpha\beta(\alpha + \beta + n)}{(\alpha + \beta)^2 (\alpha + \beta + 1)}.$$

بنابراین واریانس نسبت‌های نمونه‌ای بتا-دوجمله‌ای برابر با

$$\begin{aligned} Var\left(\frac{Z}{n}\right) &= \frac{1}{n^2} Var(Z) \\ &= \frac{\alpha\beta(\alpha + \beta + n)}{n(\alpha + \beta)^2 (\alpha + \beta + 1)} \end{aligned}$$

خواهد بود. از طرفی $Cov(Z_i, Z_j|Y) = 0$ ، بنابراین ازاره فضایی نسبت‌های نمونه‌ای، باید نسبت به ازاره توزیع بتا تفاوت چندانی نداشته باشد، یعنی $\sigma_Y^2 = Var(Y)$ باشد. البته خطای برآورد ناچیزی در آن وجود دارد. تغییرپذیری اضافی از ساختار بتا-دوجمله‌ای می‌تواند به عنوان اثر قطعه‌ای در نظر گرفته شود، یعنی $\tau^2 = Var\left(\frac{Z}{n}\right) - Var(Y)$. همچنین دامنه فضایی در سراسر تبدیل نباید تغییر چندانی داشته باشد، چون همان‌طور که عنوان شد تبدیل از میدان تصادفی گاوسی به میدان تصادفی بتای یک به یک است. همچنین از مکان‌های نمونه‌ای مشابه در هر نقطه استفاده شده است. البته ممکن است نوسان ناچیزی در برآورد دامنه فضایی در سطح نمونه‌گیری دوجمله‌ای رخ دهد، اما به طور کلی باید ثابت باشد. بنابراین پارامترهای تغییرنگار برای نسبت‌های نمونه‌ای طبق پارامترهای شکل برابر هستند با

$$\sigma_Z^2 = \sigma_Y^2 = \frac{\alpha\beta}{(\alpha + \beta)^2 (\alpha + \beta + 1)}$$

$$\phi_Z = \phi_Y$$

$$\begin{aligned}
\tau_Z^2 &= Var\left(\frac{Z}{n}\right) - Var(Y) \\
&= \frac{\alpha\beta(\alpha+\beta+n)}{n(\alpha+\beta)^2(\alpha+\beta+1)} - \frac{n\alpha\beta}{n(\alpha+\beta)^2(\alpha+\beta+1)} \\
&= \frac{\alpha\beta(\alpha+\beta)}{n(\alpha+\beta)^2(\alpha+\beta+1)} \\
&= \frac{\alpha\beta}{n(\alpha+\beta)(\alpha+\beta+1)}.
\end{aligned}$$

با توجه به این که این فرمول اندازه نمونه نابرابر را نشان نمی‌دهد، با این وجود نکته بسیار مهمی را در مورد اندازه نسبی نسبت‌های نمونه‌ای بتا-دوجمله‌ای بیان می‌کند. اثر قطعه‌ای τ_Z^2 توسط عامل $\frac{1}{n}$ تغییر می‌کند، یعنی با افزایش حجم نمونه در هر مکان فضایی، اثر قطعه‌ای به سرعت کاهش می‌یابد. بنابراین تصحیح ارایه‌شده توسط کریگینگ بتا-دوجمله‌ای زمانی بیش‌تر مورد توجه قرار دارد که تعداد مشاهدات نسبتاً کوچکی در هر مکان فضایی وجود داشته باشد. در حالت کلی می‌توان واریانس مجموعه داده‌های فضایی را در دو مولفه اثر قطعه‌ای و ازاره فضایی تجزیه کرد. یعنی

$$Var\left(\frac{Z}{n}\right) = \tau_Z^2 + \sigma_Z^2$$

که مقادیر آن برای توزیع‌های مختلف بتا در جدول ۶.۳ آورده شده‌اند. همچنین جدول ۷.۳ مقادیر ازاره و اثر قطعه‌ای را برای توزیع‌های مختلف به‌ازای n های متفاوت نشان می‌دهد. با توجه به مقادیر جدول، با افزایش n اثر قطعه‌ای به سرعت کاهش می‌یابد و از بین می‌رود. این کاهش در شکل ۴.۳ نیز با افزایش اندازه نمونه برای توزیع‌های مختلف بتا مشهود است. در نمونه‌های نسبتاً بزرگ، تغییرنگار تجربی نسبت‌های نمونه‌ای باید عملکرد خوبی داشته باشند. عبارت قطعه‌ای اضافی موجود به‌عنوان کسری از تغییرپذیری کلی، نسبتاً کوچک خواهد بود. واضح است که برای نمونه‌های کوچک استفاده از نسبت‌های نمونه‌ای برای برآورد تغییرنگار میدان بتای پنهان مناسب نخواهد بود.

جدول ۶.۳: واریانس واقعی نسبت‌های نمونه‌ای دوجمله‌ای تحت توزیع‌های بتا در شرایط شبیه‌سازی با $n=3$

موارد	پارامترهای بتا	τ_Z^2	σ_Z^2	$Var\left(\frac{Z}{n}\right)$
۱	$Beta(0.5, 0.5)$	۰/۰۴۲	۰/۱۲۵	۰/۱۶۷
۲	$Beta(1, 1)$	۰/۰۵۶	۰/۰۸۳	۰/۱۳۹
۳	$Beta(2, 2)$	۰/۰۶۷	۰/۰۵۰	۰/۱۱۷
۴	$Beta(5, 1)$	۰/۰۴۰	۰/۰۲۰	۰/۰۶۰
۵	$Beta(1, 3)$	۰/۰۵۰	۰/۰۳۸	۰/۰۸۸
۶	$Beta(2, 5)$	۰/۰۶۰	۰/۰۲۶	۰/۰۸۵

جدول ۷.۳: پارامترهای فضایی واقعی برای نسبت نمونه‌ای $\frac{Z}{n}$ با اندازه نمونه برابر

توزیع بتا	(۰/۵, ۰/۵)	(۱, ۱)	(۲, ۲)	(۵, ۱)	(۱, ۳)	(۲, ۵)
σ_Z^2	۰/۱۲۵	۰/۰۸۳	۰/۰۵۰	۰/۰۲۰	۰/۰۳۸	۰/۰۲۶
$\tau_Z^2(n=1)$	۰/۱۲۵	۰/۱۶۷	۰/۰۲۰۰	۰/۱۱۹	۰/۱۵۰	۰/۱۷۹
$\tau_Z^2(n=2)$	۰/۰۶۳	۰/۰۸۳	۰/۱۰۰	۰/۰۶۰	۰/۰۷۵	۰/۰۸۹
$\tau_Z^2(n=5)$	۰/۰۲۵	۰/۰۳۳	۰/۰۴۰	۰/۰۲۴	۰/۰۳۰	۰/۰۳۶
$\tau_Z^2(n=10)$	۰/۰۱۳	۰/۰۱۷	۰/۰۲۰	۰/۰۱۲	۰/۰۱۵	۰/۰۱۸
$\tau_Z^2(n=20)$	۰/۰۰۶	۰/۰۰۸	۰/۰۱۰	۰/۰۰۶	۰/۰۰۸	۰/۰۰۹

۲.۳ روش‌های مختلف برآورد تغییرنگار

در متون آمار فضایی (کرسی، ۱۹۸۵)، برای برآورد تغییرنگار سه استراتژی معمول برای تعیین تعداد نقاط همسایگی و تابع زیان به صورت زیر پیشنهاد شده‌اند:

۱.

$$N(h) = \sum_{i,j} I_d(i,j) \sim h \quad L_1(\theta) = \sum_k N(h) [\gamma_Y(k) - \gamma_\theta(k)]$$

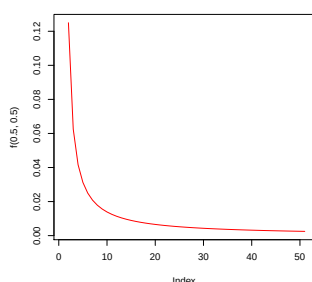
۲.

$$N^*(h) = \sum_{i,j} \frac{n_i n_j}{n_i + n_j} I_d(i,j) \sim h \quad L_2(\theta) = \sum_k N^*(h) [\gamma_Y(k) - \gamma_\theta(k)]$$

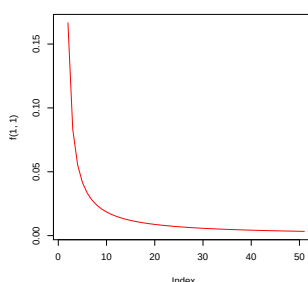
۳.

$$N(h) = \sum_{i,j} I_d(i,j) \sim h \quad L_3(\theta) = \sum_k \left[\frac{\gamma_Y^*(k) - \gamma_\theta(k)}{\gamma_\theta(k)} \right]^2$$

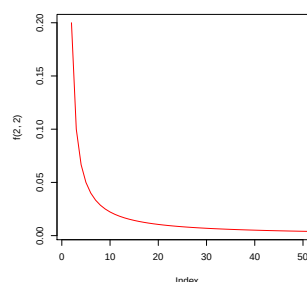
به‌طوری که در آن θ بردار پارامترهای تغییرنگار و $\gamma_Y^*(k)$ مقدار تغییرنگار بتا-دوجمله‌ای آن‌ها برای فاصله k است. توابع زیان استراتژی اول و دوم یکسان هستند و تنها اختلاف این دو استراتژی در تعداد مشاهدات در هر فاصله است. استراتژی اول فاقد وزن است و از تعداد جفت مشاهدات در هر کلاس فاصله استفاده می‌کند. استراتژی دوم از وزن تعداد مشاهدات یعنی $\frac{n_i n_j}{n_i + n_j}$ برای (i, j) استفاده می‌کند. در حقیقت این عبارت وزن‌دار، اطلاعات مکان‌های مختلف را همگن می‌کند و به جفت مکان‌ها با اندازه نمونه بیش‌تر، وزن بیش‌تری را اختصاص می‌دهد. استراتژی سوم نیز فاقد وزن است.



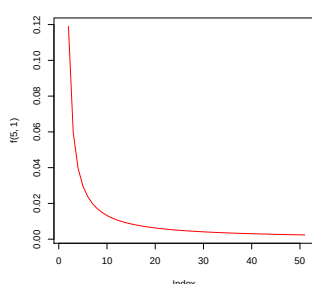
$Beta(0.5, 0.9)$ (ج)



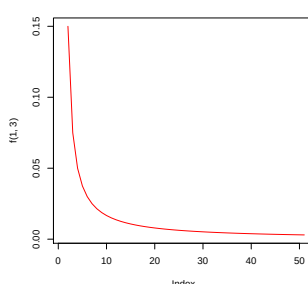
$Beta(1, 1)$ (ب)



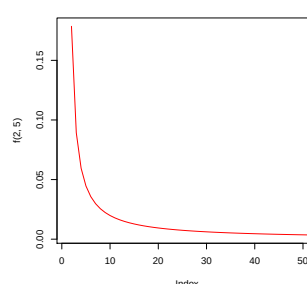
$Beta(2, 2)$ (آ)



$Beta(5, 1)$ (و)



$Beta(1, 3)$ (ه)

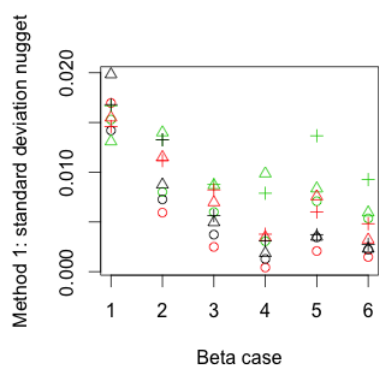


$Beta(2, 5)$ (د)

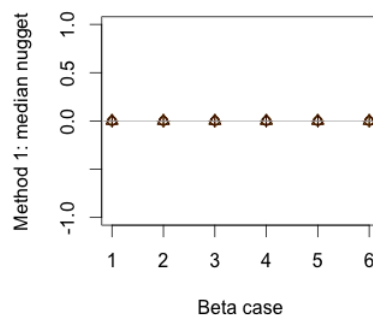
شکل ۴.۳: اثر قطعه‌ای نظری توزیع‌های بتا

برای هر یک از ۵۴ سناریو شامل شش توزیع بتا، سه اندازه نمونه مختلف و سه دامنه فضایی، ۱۰۰ مجموعه داده شبیه‌سازی کردیم. سپس برای تمام مجموعه‌های تولیدشده هر سه استراتژی برآورد پارامتر را اجرا کردیم. شکل‌های ۵.۳، ۶.۳ و ۷.۳ برآورد میانه و انحراف معیار را برای ۵۴ سناریو با استراتژی‌های مختلف نشان می‌دهند. در این شکل‌ها دامنه ۵ با $\circ$ ، 10 با $\triangle$ و 15 با $+$ نشان داده شده‌اند. همچنین اندازه نمونه کوچک با رنگ سبز، اندازه نمونه متوسط با رنگ قرمز و اندازه نمونه بزرگ با رنگ سیاه نمایش داده شده‌اند. با توجه به نتایج شبیه‌سازی‌ها، استراتژی سوم نسبت به دو استراتژی دیگر بدترین عملکرد را نشان می‌دهد. در این روش اثر قطعه‌ای برای توزیع‌های متقارن به‌خوبی برآورد شده است، اما تقریباً برای تمام توزیع‌های نامتقارن بیش‌برآورد شده است. انحراف معیار اثر قطعه‌ای استراتژی سوم نسبت به سایر روش‌ها بزرگ‌تر است. از طرفی در این استراتژی بیش‌ترین اثر قطعه‌ای در اندازه نمونه کوچک وجود دارد، در حالی که دو استراتژی دیگر اثر قطعه‌ای را برای هر سه اندازه نمونه به‌خوبی برآورد می‌کنند. همچنین در این استراتژی ازاره فضایی برای تمام سناریوها بیش‌تر برآورد شده است. ازاره فضایی برای توزیع‌های متقارن با کاهش دامنه کم‌تر برآورد می‌شود. کرسی (۱۹۹۳) استفاده از استراتژی سوم را برای توزیع‌های گاوسی پیشنهاد داد؛ بنابراین عملکرد نامناسب این روش برای داده‌های دوجمله‌ای شمارشی کاملاً طبیعی است.

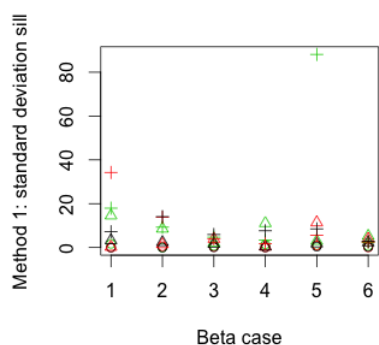
حال به مقایسه دقت برآوردهای استراتژی اول و دوم می‌پردازیم. هر دو استراتژی اثر قطعه‌ای را برای تمام سناریوها به‌خوبی برآورد می‌کنند. به‌طور متوسط می‌توان گفت انحراف معیار اثر قطعه‌ای برای روش دوم تا اندازه‌ای کوچک‌تر است. در روش دوم دو سناریو با انحراف معیار بالا وجود دارند که ممکن است به علت وجود چولگی در داده‌ها رخ داده باشند. برآورد میانه ازاره، تقریباً در هر دو روش یکسان است. استراتژی دوم دامنه فضایی را کمی بهتر برآورد می‌کند. روش اول در دامنه فضایی 10° توزیع‌های ۵-۶ را کمی بیش‌تر و همچنین ۱-۳ را کم‌تر برآورد می‌کند. در استراتژی اول مشکل بیش‌برآوردی برای توزیع‌های متقارن وجود دارد، اما برای توزیع‌های چوله دقیق است. همچنین این استراتژی دامنه فضایی ۱۵ را برای توزیع‌های متقارن کم‌برآورد می‌کند و برای توزیع‌های چوله کاملاً متغیر است. روش دوم نیز توزیع بتا ۱ و ۲ را بیش‌برآورد می‌کند، اما برای سایر سناریوهای شبیه‌سازی به‌ویژه در اندازه نمونه بزرگ دقیق است. بر اساس شبیه‌سازی‌ها، استراتژی دوم برآوردهای دقیق‌تری از پارامترهای کوواریانس را ارائه می‌دهد. تعدیل برای تعداد مشاهدات در هر فاصله h یعنی $N^*(h)$ به‌صورت تصادفی انتخاب نشده است. در حقیقت این همان تعدیل ارائه‌شده در کریگینگ بتا-دوجمله‌ای است. این تعدیل نشان می‌دهد که همگن کردن اندازه نمونه n_i در مکان‌های مختلف برای برآورد پارامتر نیز مهم است. هنگام برخورد با نسبت‌ها، دقت نسبت‌های دوجمله‌ای به‌ویژه در اندازه نمونه کوچک به اندازه نمونه وابسته است. این واقعیت نه تنها در مرحله برآورد تغییرنگار بلکه در مرحله برآورد پارامترهای کوواریانس مدل نیز باید مورد توجه قرار گیرد.



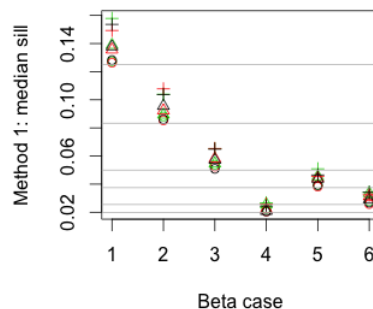
(ب) انحراف معیار اثر قطعه‌ای



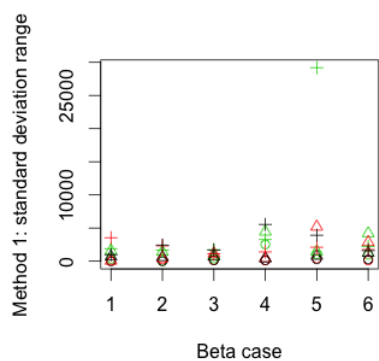
(آ) میانه اثر قطعه‌ای



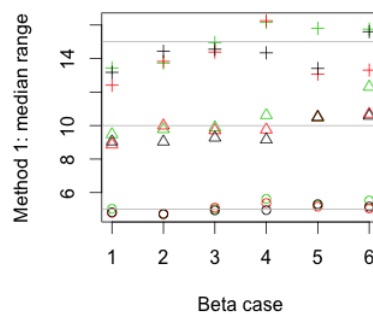
(د) انحراف معیار ازاره



(ج) میانه ازاره

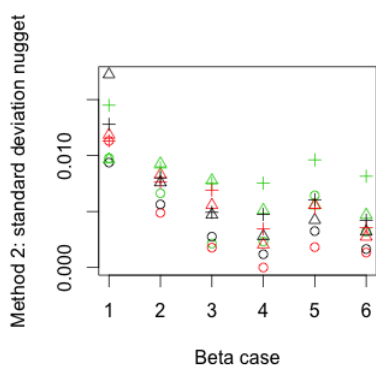


(و) انحراف معیار دامنه

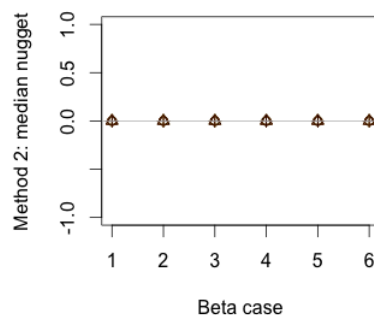


(ه) میانه دامنه

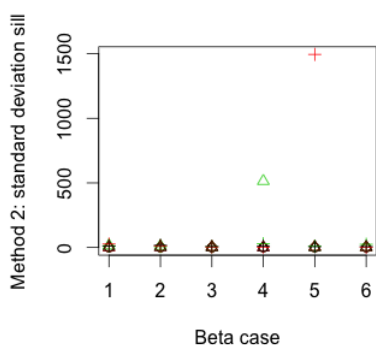
شکل ۵.۳: برآورد میانه و انحراف معیار پارامترهای تغییرنگار برای ۵۴ سناریو شبیه‌سازی بر اساس ۱۰۰ مجموعه داده شبیه‌سازی با استفاده از روش اول



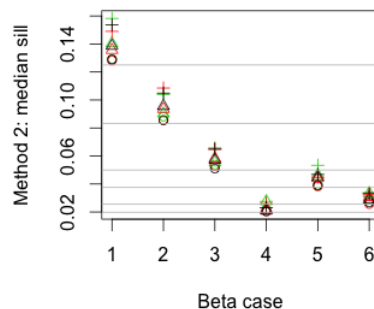
(ب) انحراف معیار اثر قطعه‌ای



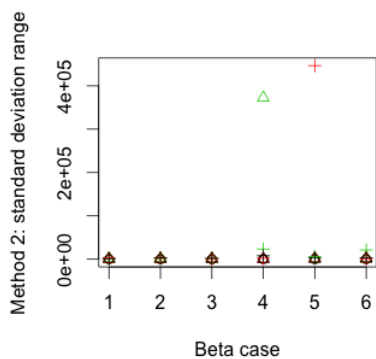
(آ) میانه اثر قطعه‌ای



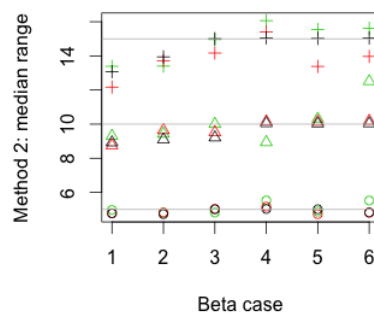
(د) انحراف معیار ازاره



(ج) میانه ازاره

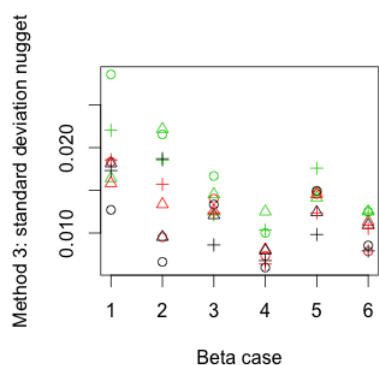


(و) انحراف معیار دامنه

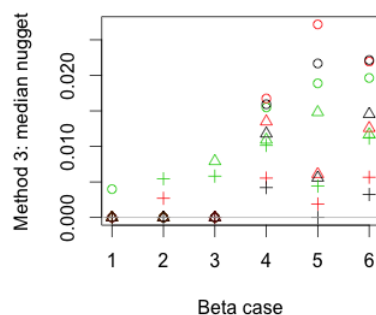


(ه) میانه دامنه

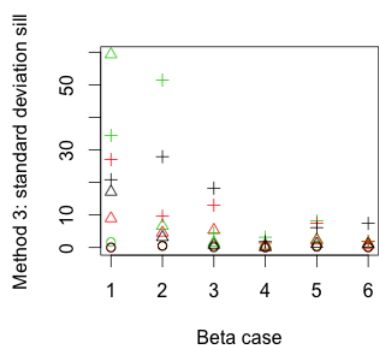
شکل ۶.۳: برآورد میانه و انحراف معیار پارامترهای تغییرنگار برای ۵۴ سناریو شبیه‌سازی بر اساس ۱۰۰ مجموعه داده شبیه‌سازی با استفاده از روش دوم



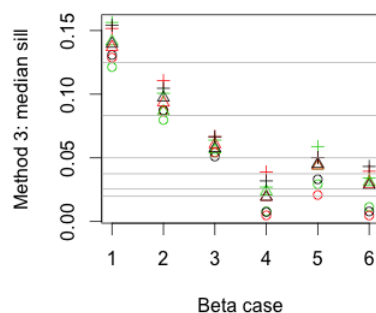
(ب) انحراف معیار اثر قطعه‌ای



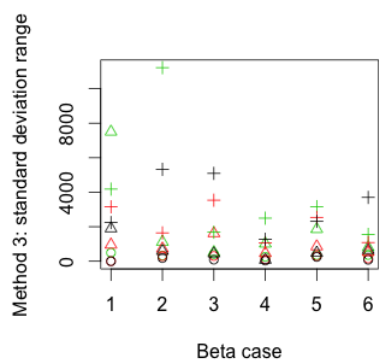
(آ) میانه اثر قطعه‌ای



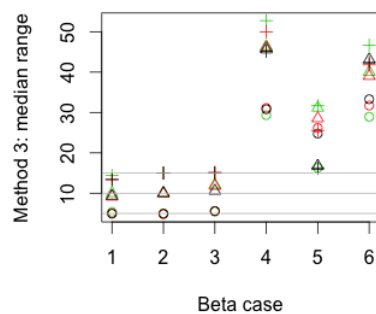
(د) انحراف معیار ازاره



(ج) میانه ازاره



(و) انحراف معیار دامنه



(ه) میانه دامنه

شکل ۷.۳: برآورد میانه و انحراف معیار پارامترهای تغییرنگار برای ۵۴ سناریو شبیه‌سازی بر اساس ۱۰۰ مجموعه داده شبیه‌سازی با استفاده از روش سوم

۳.۳ برآورد پارامترهای ساختار وابستگی فضایی

برای ارزیابی عملکرد برآورد پارامترها، برای هر سناریو ۱۰۰ مجموعه داده شبیه‌سازی کردیم. همان‌طور که در بخش قبلی بیان شد، با توجه به نتایج حاصل از شبیه‌سازی، برای برآورد پارامترها از تابع زیان $L_2(\cdot)$ و تعریف تعداد نقاط همسایه $N^*(h)$ ، در استراتژی دوم، استفاده کردیم. برآوردهای پارامترهای تغییرنگار برای مدل بتا-دوجمله‌ای از می‌نیم کردن تابع زیان $L_2(\theta)$ ، که در تغییرنگار تجربی و تغییرنگار واقعی (کروی) مقداردهی شده است، به دست آمده‌اند. نمودارهای جعبه‌ای برآوردهای پارامترهای اثر قطعه‌ای، ازاره و دامنه برای مدل بتا-دوجمله‌ای در شکل‌های ۸.۳، ۹.۳ و ۱۰.۳ نشان داده شده‌اند. در این شکل، ۱۸ سناریوی اول مربوط به دامنه ۵، ۱۸ سناریو دوم مربوط به دامنه ۱۰ و ۱۸ سناریو سوم مربوط به دامنه ۱۵ هستند. همچنین ۶ سناریو اول هر دامنه فضایی، اندازه نمونه کوچک، ۶ سناریو دوم اندازه نمونه متوسط و ۶ سناریو سوم اندازه نمونه بزرگ را نشان می‌دهند. پارامترهای واقعی دامنه فضایی به ترتیب برای این سناریوها ۵، ۱۰ و ۱۵ می‌باشند. اثر قطعه‌ای واقعی صفر است که تقریباً در تمام شبیه‌سازی‌ها به درستی برآورد شده است. پارامترها را برای هر یک از ۱۰۰ مجموعه داده شبیه‌سازی‌شده با استفاده از استراتژی دوم که در بخش قبل مطرح شد برآورد کردیم. در این بخش مقادیر برآوردشده را برای سه پارامتر فضایی اثر قطعه‌ای، دامنه و ازاره مورد بررسی قرار می‌دهیم.

پارامتر اثر قطعه‌ای یعنی τ^2 تغییرات ایجادشده توسط نمونه‌های دوجمله‌ای در هر مکان را نشان می‌دهد. در تمام شبیه‌سازی‌ها تغییرپذیری اثر قطعه‌ای برای میدان تصادفی بتا پنهان صفر است، بنابراین هرگونه تغییرپذیری اثر قطعه‌ای در نسبت‌های نمونه‌ای فقط از نمونه‌های دوجمله‌ای ناشی شده است و احتمال زیربنایی در آن تاثیری ندارد. با توجه به این‌که پارامترها را برای ساختار فضایی فرآیند بتای پنهان برآورد می‌کنیم، بنابراین در تمام شبیه‌سازی‌ها میانه اثر قطعه‌ای صفر خواهد بود، که نشان می‌دهد اصلاح برآوردگر تغییرنگار در از بین بردن تنوع اضافی موفق بوده است. با این وجود در برخی از توزیع‌های بتا هنوز هم کمی تغییرپذیری موجود است. به‌عنوان مثال، توزیع بتا نوع یک که نشان‌دهنده توزیع بتا دومدی است بیش‌ترین اثر قطعه‌ای برآوردشده را دارد. به‌طور کلی این توزیع بیش‌ترین تنوع را در داده‌های نمونه‌ای نشان داده است. بنابراین باید انتظار داشت که تنوع اضافی همواره به‌طور کامل نمی‌تواند تعدیل شود. توزیع‌های بتا نامتقارن (۴، ۵، ۶) تمایل به برآورد نقاط پرت کم‌تری دارند. انحراف معیار اثر قطعه‌ای برآوردشده نیز برای توزیع‌های متقارن تمایل دارند افزایش پیدا کنند. از آنجایی که اثر قطعه‌ای در صفر کراندار است، افزایش انحراف معیار تنها به‌دلیل اثر قطعه‌ای مثبت بزرگ است.

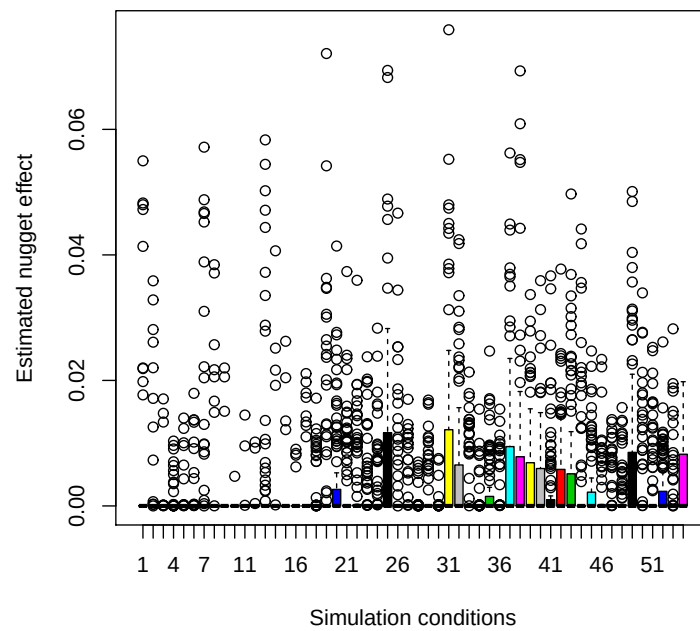
همچنین به‌نظر می‌رسد اندازه اثر قطعه‌ای برآوردشده به دامنه فضایی وابسته است. ۱۸ سناریو اول که دارای کم‌ترین دامنه فضایی، یعنی $\phi = 5$ ، هستند نقاط پرت نسبتاً کم‌تری

را دارند. سناریوهای ۵۴-۱۹ ازاره فضایی بزرگتر و نقاط پرت بهمراتب بیشتری را نشان می‌دهند. از طرفی با افزایش دامنه فضایی، مقادیر نقاط پرت تمایل دارند افزایش پیدا کنند. می‌توان گفت اندازه نمونه در برآورد اثر قطعه‌ای تاثیری ندارد. در واقع با افزایش اندازه نمونه الگوی قابل تشخیصی در شکل معلوم نیست، که این نشان می‌دهد اصلاح تغییرنگار بتا-دوجمله‌ای، برآورد پارامترهای فرآیند بتای پنهان را بهبود بخشیده است. با توجه به این که در اندازه نمونه کوچک اطلاعات کمتری درباره نسبت بتا درست وجود دارد، بنابراین با استفاده از تغییرنگار اصلاح نشده در این اندازه نمونه باید نتایج تغییرپذیری بیشتری داشته باشند.

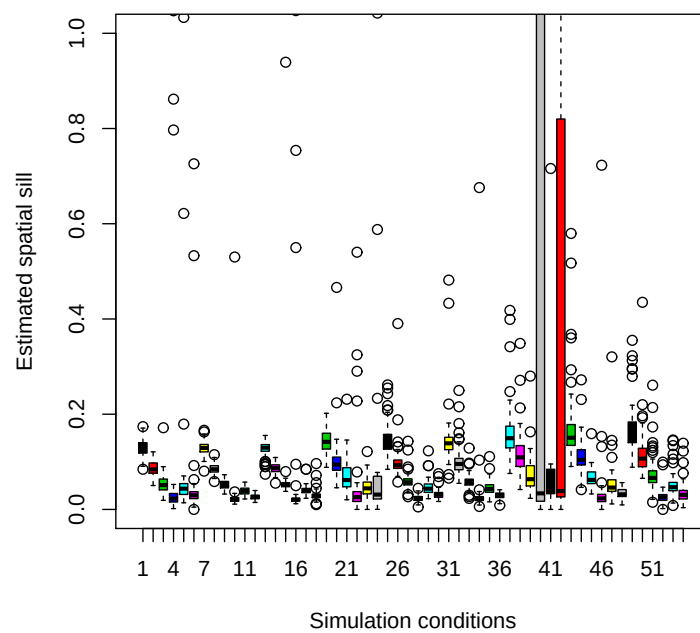
ازاره فضایی مورد انتظار توزیع بتای پنهان، تغییرپذیرتر از توزیع بتای زیربنایی است. نمودار جعبه‌ای برآورد پارامتر ازاره، چوله به راست است و نقاط پرت کمتری نسبت به اثر قطعه‌ای دارد. از طرفی، برآوردهای ازاره بزرگتر از یک نیز وجود دارند که با توجه به کراندار بودن توزیع بتا در فاصله صفر و یک طبیعی نیست و نشان می‌دهد که همگرایی برای بعضی از توزیع‌ها ممکن است مشکل ساز شود. همچنین با افزایش دامنه فضایی نقاط پرت نیز افزایش می‌یابند.

اندازه نمونه‌های دوجمله‌ای روی میانه ازاره فضایی برای توزیع‌های متقارن تاثیر ندارند. با افزایش دامنه فضایی، نمودار میانه ازاره، ازاره فضایی واقعی را بیش‌تر برآورد می‌کند، به‌طوری که در اندازه نمونه بزرگ این بیش‌برآوردی کاهش می‌یابد اما هنوز هم قابل مشاهده است. تغییرپذیری اضافی می‌تواند به عوامل مختلفی وابسته باشد. یک عامل ممکن می‌تواند برآورد پارامترهای توزیع بتا با استفاده از نسبت‌های نمونه‌ای باشد. با استفاده از $\hat{\alpha}$ و $\hat{\beta}$ برآورد شده برای برآورد تغییرنگار بتای پنهان، سطح دیگری از پیچیدگی اصلاح نشده در تغییرنگار اضافه می‌شود که می‌تواند منجر به متورم شدن اثر واریانس شود. توزیع‌های بتا متقارن دامنه گسترده‌تری از مقادیر کلی ممکن دارند و از توزیع‌های بتا چوله متغیرتر هستند.

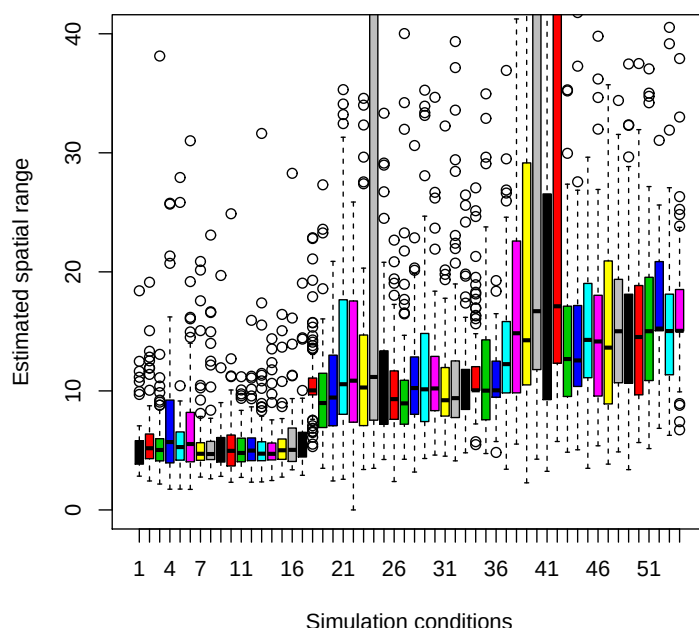
در نهایت دامنه فضایی برآورد شده را در نظر بگیرید. ۱۸ سناریو اول مربوط به کوچک‌ترین دامنه فضایی، یعنی $\phi = 5$ است. به‌طور کلی برآورد دامنه فضایی در شبیه‌سازی‌ها به‌خوبی بازیابی شده است. اگرچه هنوز هم چولگی کمی در برآورد پارامترها وجود دارد که با افزایش اندازه نمونه کاهش می‌یابد. بر اساس میدان تصادفی فضایی با ابعاد 20×20 ، دامنه فضایی ۵ با توجه به این که فاصله نسبتاً کمی با آن دارد تشخیص آن باید آسان باشد. مدل کریگینگ بتا-دوجمله‌ای، برآورد دامنه فضایی در این اندازه را به‌خوبی انجام می‌دهد. همچنین تغییرپذیری توزیع‌های چوله نیز بیش‌تر از توزیع‌های بتا متقارن است. الگوهای افزایش برآورد دامنه برای توزیع‌های چوله با دامنه $\phi = 10$ و $\phi = 15$ مشابه دامنه $\phi = 5$ رخ می‌دهد. به‌طور کلی می‌توان گفت، مدل کریگینگ بتا-دوجمله‌ای برآورد پارامترها را به درستی انجام می‌دهد.



شکل ۸.۳: اثر قطعه‌ای فضایی برآورد شده



شکل ۹.۳: ازاره فضایی برآورد شده



شکل ۳.۱۰: دامنه فضایی برآورد شده

برآورد میدان تصادفی بتای پنهان، دارای معایبی است که باید مرتفع شوند. این معایب شامل موارد زیر هستند:

۱. اطلاعات محدود هستند. در اغلب موارد به‌ویژه در اندازه نمونه کوچک نسبت‌های نمونه‌ای دوجمله‌ای می‌توانند بیش‌تر یا کم‌تر از احتمال زیربنایی باشند.

۲. پارامترها پیچیده‌تر هستند. نه تنها پارامترهای کوواریانس فضایی باید برآورد شوند، بلکه ویژگی‌های میدان تصادفی بتای زیربنایی نیز باید به‌خوبی برآورد شوند. در تحلیل زمین‌آماري، فرض نرمال بودن میدان تصادفی فضایی و ثابت بودن میانگین برای برآورد پارامترهای ساختار وابستگی فضایی ضروری است.

۳. اندازه نمونه ثابت نیست. در هر نقطه سطح متفاوتی از تغییرپذیری در نسبت‌های نمونه‌ای دوجمله‌ای وجود دارد، که این تنوع باید برای عملکرد میدان تصادفی بتای پنهان برآورد و سپس حذف شود.

مدل کریگینگ بتا-دوجمله‌ای به‌عنوان رهیافتی امیدبخش برای مدل‌سازی مسایل پیچیده ظاهر شده است. این مدل برآورد دقیق و پایداری را از ساختار فضایی میدان احتمالی پنهان فراهم می‌کند.

۴.۳ دقت پیش‌گویی

برای بررسی عملکرد مدل بتا-دوجمله‌ای در پیش‌گویی، ۱۰۰ مجموعه داده را (به ازای هر سناریوی شبیه‌سازی) بر روی یک شبکه 50×50 شامل ۲۵۰۰ گره تولید کردیم و با استفاده از روش کریگینگ بتا-دوجمله‌ای (با پارامترهای برآوردشده تغییرنگار) مقادیر نسبت‌های دوجمله‌ای را، برای هر مجموعه داده، پیش‌گویی کردیم. در واقع در این جا می‌خواهیم بررسی کنیم که آیا مدل، مقادیر مشاهده‌نشده فرآیند بتای پنهان را به‌خوبی پیش‌گویی کرده است یا خیر؟ برای این منظور از معیارهای مختلفی استفاده می‌کنیم، که در ادامه به تحلیل نتایج این معیارها می‌پردازیم.

نمودارهای شکل‌های ۱۱.۳، ۱۲.۳، ۱۳.۳، ۱۴.۳، ۱۵.۳ و ۱۶.۳ معیارهای مختلف دقت پیش‌گویی را نمایش می‌دهند. در شکل ۱۱.۳ نمودار جعبه‌ای میانگین خطای پیش‌گویی شده را برای هر یک از ۱۰۰ مجموعه داده شبیه‌سازی شده تحت سناریوهای مختلف نشان می‌دهد. بر اساس همه سناریوهای شبیه‌سازی، میانگین خطای پیش‌گویی بسیار نزدیک به صفر است و روند کمی در آن مشاهده می‌شود. به‌طور کلی با افزایش اندازه نمونه، تغییرات میانگین خطای پیش‌گویی کوچک و تقریباً صفر است. همچنین در دامنه فضایی بزرگ‌تر نیز خطای پیش‌گویی نزدیک به صفر است. از طرفی تعدادی نقاط پرت دیده می‌شوند که نشان می‌دهد حتی اگر برآورد پارامترها برای یک مجموعه داده دشوار باشد مقادیر پیش‌گویی منطقی و دقیق هستند. توزیع بتا ۱ و ۲ که نشان‌دهنده توزیع بتا دومدی و توزیع یکنواخت هستند، به‌ترتیب بیش‌ترین تغییرات میانگین خطای پیش‌گویی را نشان می‌دهند. این توزیع‌های خاص بیش‌تر تمایل دارند در نزدیکی مقادیر صفر و یک که پیش‌بینی با استفاده از مدل کریگینگ بتا-دوجمله‌ای دشوار است، رخ دهند.

میانگین توان دوم خطای پیش‌گویی یک شاخص خوب از دو ویژگی مجموعه‌ای از پیش‌گویی‌ها، یعنی دقت و تنوع است. مجموعه داده‌های پیش‌گویی دقیق باید میانگین توان دوم خطای پیش‌گویی نزدیک به صفر داشته باشند. به‌طوری که افزایش میانگین توان دوم خطای پیش‌گویی در دو حالت رخ می‌دهد: این افزایش می‌تواند به دلیل خطای نسبتاً بالا یا وجود نقاط پرت در خطاها باشد. توزیع‌های بتا متقارن در تمام شبیه‌سازی‌ها بیش‌ترین میانگین توان دوم خطای پیش‌گویی را دارند، که علت آن متغیرتر بودن این توزیع‌ها است. همچنین این توزیع‌ها تمایل دارند نسبت‌های شدیدتر را نمایش دهند. توزیع‌های چوله نیز در تمامی شبیه‌سازی‌ها میانگین توان دوم خطای پیش‌گویی کوچک و تقریباً کم‌تر از ۰.۰۴ را دارند. علاوه بر این افزایش اندازه نمونه سبب کاهش میانگین توان دوم خطای پیش‌گویی می‌شود. بنابراین می‌توان گفت، با گرفتن نمونه‌های دوجمله‌ای، پیش‌گویی‌ها تمایل دارند دقیق‌تر شوند که این کاهش بیش‌تر از نمونه‌های کوچک به متوسط مشهود است. این نشان می‌دهد که با مدل کریگینگ بتا-دوجمله‌ای، ده تا بیست مشاهده در هر مکان فضایی ممکن است برای دریافت نتیجه معقول با اطلاعات خوب کافی باشد. با افزایش دامنه فضایی، میانگین توان

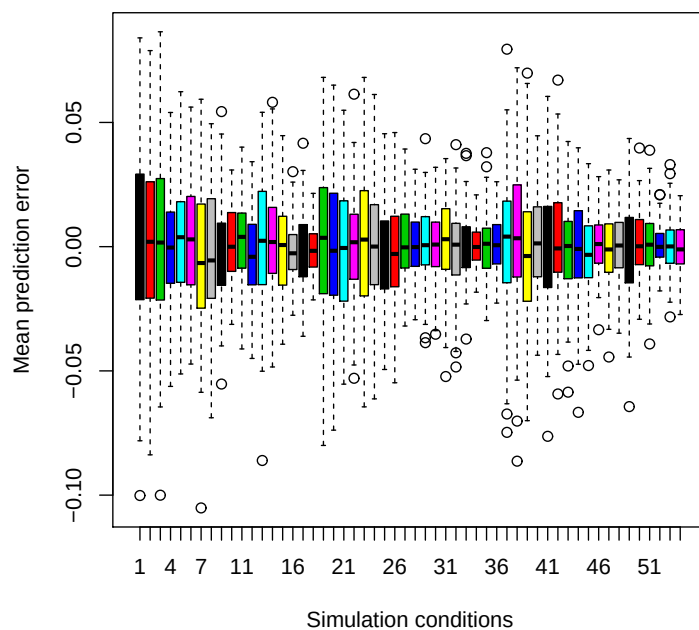
دوم خطای پیش‌گویی کاهش می‌یابد. هرچه وابستگی فضایی قوی‌تر باشد، تغییرات از مکانی به مکان دیگر کم‌تر، و پیش‌گویی نیز می‌تواند بهتر باشد.

معیار دیگر دقت پیش‌گویی، میانگین مقادیر پیش‌گویی است. آیا میانگین مقادیر پیش‌گویی به میانگین واقعی توزیع بتا نزدیک هستند؟ با استفاده از مدل کریگینگ بتا-دوجمله‌ای پاسخ مثبت است، به‌طوری که میانگین مقادیر پیش‌گویی برای شبیه‌سازی‌های بتا به میانگین توزیع بتا بسیار نزدیک هستند. در اندازه نمونه بزرگ میانگین مقادیر پیش‌گویی بسیار سازگار هستند. همان‌طور که انتظار داشتیم توزیع‌هایی با دامنه فضایی بزرگ، مقادیر میانگین پیش‌گویی بسیار پایداری دارند.

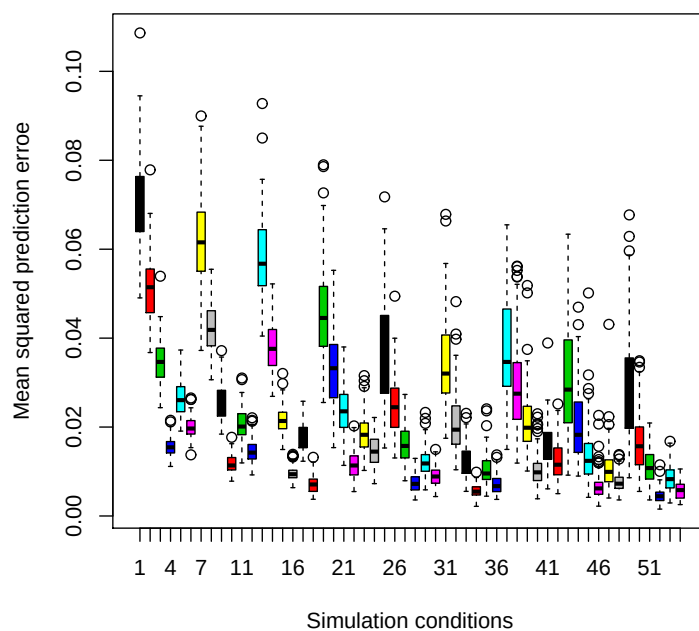
با توجه به شکل ۱۶.۳ مجموعه داده‌ها با دامنه فضایی کوچک مقادیر پیش‌گویی با واریانس کوچک‌تری دارند. به‌طور کلی نمای پیش‌گویی‌ها مسطح است، که این مسطح بودن به خطاهای بالاتر مشاهده‌شده در این شبیه‌سازی برمی‌گردد. سطح پیش‌گویی‌هایی که صاف هستند، گوشه‌ها و شکاف داده‌ها را نمی‌توانند به‌خوبی برآورد کنند. در دامنه فضایی بزرگ‌تر سطوح پیش‌گویی شیب‌دار و واریانس بزرگ‌تر می‌شود. همچنین واریانس پیش‌گویی برای توزیع‌های متقارن بزرگ‌تر است، همان‌طوری که توزیع بتا زیربنایی متغیرتر است. بنابراین پیش‌گویی‌ها نیز باید متغیرتر باشند. با افزایش اندازه نمونه واریانس پیش‌گویی نیز افزایش می‌یابد.

با معیارهای پیش‌گویی که تاکنون در نظر گرفتیم، کاهش برای اندازه نمونه بزرگ به‌خوبی صورت گرفته است. با این وجود، برای واریانس پیش‌گویی در اندازه نمونه بزرگ‌تر باید امیدوار به پیش‌گویی‌های متغیرتر بود. نمونه‌های دوجمله‌ای با n_i بزرگ تمایل دارند به نسبت‌های نمونه‌ای مشاهده‌شده با احتمال درست π_i نزدیک‌تر شوند. همان‌طور که اندازه نمونه افزایش می‌یابد، تعداد قابل توجهی از نمونه‌های مشاهده‌شده نیز دیده می‌شوند. نه تنها نسبت‌های نمونه‌ای مورد استفاده برای پیش‌گویی در اندازه نمونه بزرگ دقیق‌تر هستند بلکه آن‌ها انتخاب‌های بیش‌تری دارند. در حقیقت اطلاعات بیش‌تری قابل دستیابی است. اطلاعات کامل‌تر توسط افزایش n_i گرفته‌شده، اجازه می‌دهد که برخی از قله‌ها و دره‌ها بیش‌تر قابل توجه شوند. از این رو در اندازه نمونه بزرگ‌تر پیش‌گویی‌ها متغیرتر خواهند بود.

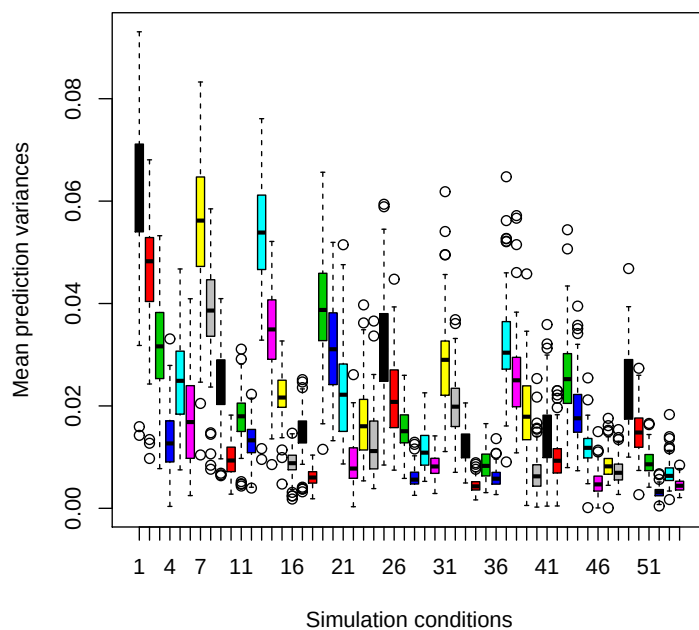
ضریب همبستگی اسپیرمن بین میدان تصادفی بتای پنهان واقعی و مقادیر پیش‌گویی را نیز برای هر شبیه‌سازی محاسبه کردیم. با افزایش اندازه نمونه، همبستگی بین میدان تصادفی بتا و متغیرهای تصادفی پیش‌گویی بیش‌تر می‌شود. ارتباط بالا به معنی برآورد بهتر بالا و پایین توزیع است. همچنین با افزایش دامنه فضایی همبستگی بین میدان توزیع بتای پنهان و توزیع بتا پیش‌گویی بیش‌تر می‌شود. این نتیجه نشان می‌دهد مدل کریگینگ بتا-دوجمله‌ای می‌تواند نقاط مرتفع و پست توزیع احتمال زیربنایی را به‌خوبی تشخیص دهد به جز در مواردی که اندازه نمونه کوچک و همبستگی فضایی نسبتاً کم باشد.



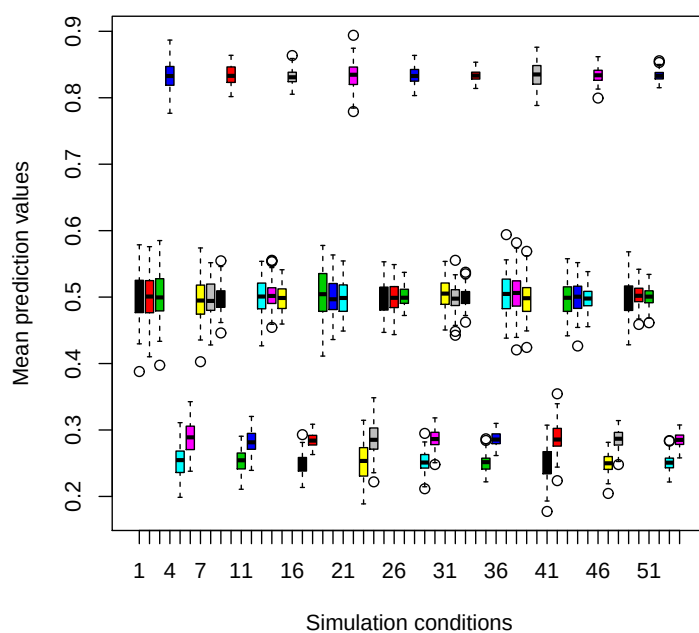
شکل ۱۱.۳: میانگین خطای پیش‌گویی



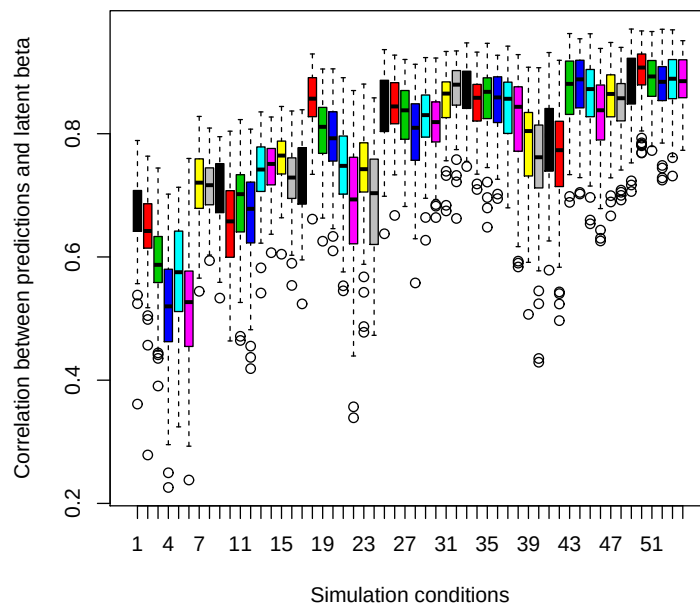
شکل ۱۲.۳: میانگین توان دوم خطای پیش‌گویی



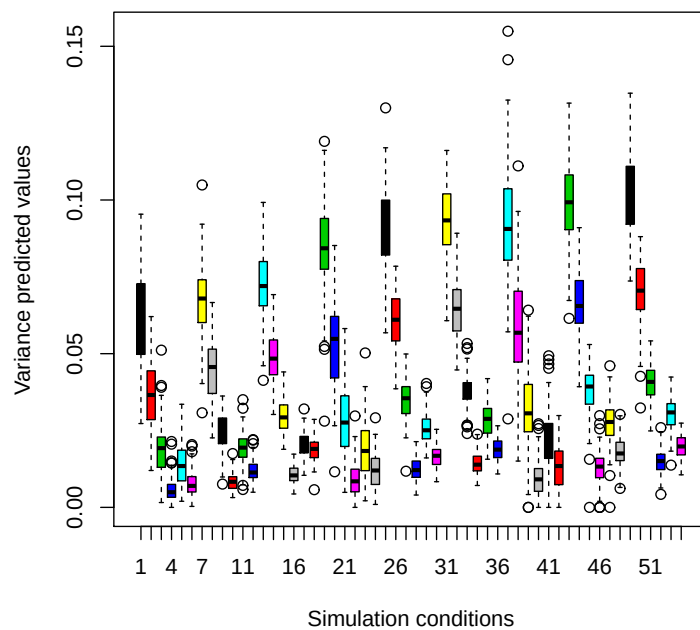
شکل ۱۳.۳: میانگین واریانس پیش‌گویی



شکل ۱۴.۳: میانگین مقادیر پیش‌گویی



شکل ۱۵.۳: همبستگی بین توزیع بتا پیش‌گویی و توزیع بتا پنهان



شکل ۱۶.۳: واریانس مقادیر پیش‌گویی

فصل ۴

کارایی پیش‌گویی مدل بتا-دوجمله‌ای فضایی و مقایسه با روش هم‌کریگینگ دوجمله‌ای

در این فصل یک مدل جایگزین برای کریگینگ بتا-دوجمله‌ای، یعنی مدل هم‌کریگینگ دوجمله‌ای را بر اساس نسبت‌های نمونه‌ای مطرح می‌کنیم. سپس به مقایسه توان هر دو مدل در برآورد پارامترها و پیش‌گویی می‌پردازیم. در انتها کاربست مدل بتا-دوجمله‌ای را برای پهنه‌بندی نرخ بارش در استان سمنان نمایش می‌دهیم.

۱.۴ هم‌کریگینگ دوجمله‌ای

لاجونی (۱۹۹۱) یک روش مشابه برای برآورد احتمال زیربنایی، بر اساس نسبت نمونه‌ای مشاهده‌شده به نام هم‌کریگینگ دوجمله‌ای^۱ معرفی نمود. در این مدل، نسبت‌های نمونه‌ای به‌عنوان متغیرهای کمکی برای میدان تصادفی پنهان مدل‌بندی می‌شوند. دلیل نام‌گذاری هم‌کریگینگ نیز همین واقعیت است. به عبار دیگر با وجود نام هم‌کریگینگ، درحقیقت

^۱Binomial cokriging

مدل به شکل هم‌کریگینگ نیست: زیرا هیچ متغیر کمکی برای کمک به پیش‌گویی فضایی استفاده نمی‌شود. در واقع بنا به گفته الیور "از آنجایی که ما یک احتمال نامعلوم را از روی نسبت‌های نمونه‌ای مشاهده‌شده برآورد می‌کنیم، بنابراین می‌توان آن را به عنوان یک مدل هم‌کریگینگ در نظر گرفت". اگرچه این روش به صورت گسترده اجرا نشد، اما در سال ۱۹۹۶ در مقاله‌ای توسط الیور و همکاران دوباره ظاهر شد. این مدل توسط لاجونی گسترش داده شد. الیور و همکاران نیز این مدل را برای پیش‌گویی خطر ابتلا به سرطان کودکان در بخش‌هایی از غرب انگلستان برای داده‌های سال‌های ۱۹۸۴-۱۹۸۰ تشریح کردند. برآوردگر تغییرنگار هم‌کریگینگ دوجمله‌ای با استفاده از نمادگذاری تغییرنگار بتا-دوجمله‌ای به صورت زیر است:

$$\hat{\gamma}_Y(h) = \hat{\gamma}_Z(h) - \frac{1}{\hat{\pi}} \left(\hat{\pi}(1 - \hat{\pi}) - \hat{\sigma}_Y^2 \right) \frac{n_i + n_j}{n_i n_j}.$$

شکل کلی برآوردگر تغییرنگار مشابه مدل کریگینگ بتا-دوجمله‌ای است. با این وجود، برخی تفاوت‌های کلیدی در هر دو تغییرنگار و همچنین پذیره‌های دو مدل وجود دارند. به طور خاص این تفاوت‌ها عبارتند از

۱. در مدل کریگینگ بتا-دوجمله‌ای به صراحت فرض می‌شود میدان تصادفی پنهان، همبسته فضایی است و از توزیع بتا پیروی می‌کند. اما در مدل هم‌کریگینگ دوجمله‌ای هیچ پذیره توزیعی بر روی مدل زیربنایی وجود ندارد.

۲. در مدل کریگینگ بتا-دوجمله‌ای یک شرط اضافی بر روی پارامترهای برآوردشده میانگین و واریانس، یعنی $\hat{\pi}$ و $\hat{\sigma}_Y^2$ ، گذاشته می‌شود تا میدان تصادفی پنهان را بین $[0, 1]$ محدود کند. در مقابل هم‌کریگینگ دوجمله‌ای چنین شرطی را لازم ندارد و به جای آن از میانگین و واریانس نسبت‌های نمونه‌ای برای برآورد $\hat{\pi}$ و $\hat{\sigma}_Y^2$ استفاده می‌شود.

۳. عبارت تصحیح اریبی در هم‌کریگینگ دوجمله‌ای مقداری ثابت است.

۴. در مدل کریگینگ بتا-دوجمله‌ای، تغییرنگار برآوردشده نسبت‌های نمونه‌ای به هر یک از جفت مکان‌ها با توجه به اختلاف واریانس‌ها وزن اختصاص می‌دهد، به طوری که با استفاده از این وزن به اختلاف بین جفت مکان‌های فضایی با اندازه نمونه بزرگ‌تر اهمیت بیشتری داده می‌شود. در مدل هم‌کریگینگ دوجمله‌ای هیچ گونه وزنی به تغییرنگار نسبت‌های نمونه‌ای اختصاص داده نمی‌شود. این عامل ممکن است سبب شود به نقاطی با اندازه نمونه کوچک اهمیت بیشتری داده شود.

با در نظر گرفتن این اختلاف‌ها، انتظار داریم مدل کریگینگ بتا-دوجمله‌ای، برآوردهای دقیق‌تری برای پارامترها و خطای پیش‌بینی کوچک‌تری تولید کند.

۱.۱.۴ معادلات هم کریگینگ دوجمله‌ای

برای این که مقدار نسبت را در مکان جدید s_o محاسبه کنیم، دو گزینه داریم. اول، تعمیم هم کریگینگ معمولی است. فرض می‌کنیم که میانگین نامعلوم است و مقدار پیش‌گویی در مکان جدید عبارت است از

$$\hat{Y}(s_o) = \sum_{i=1}^m \lambda_i \frac{Z_i}{n_i}$$

که در آن m تعداد مکان‌های مشاهده‌شده و λ_i ها وزن‌های هم کریگینگ هستند. در این جا نیز شرط ناریبی پیش‌گو $\sum_{i=1}^m \lambda_i = 1$ است. با توجه به مانایی مرتبه دوم میدان تصادفی داریم

$$E[Y(s_i)] = E\left[\frac{Z_i}{n_i}\right] = \mu.$$

بنابراین تحت شرط ناریبی و می‌نیمم کردن واریانس خطای پیش‌گویی، معادلات هم کریگینگ دوجمله‌ای به صورت

$$\hat{Y}(s_o) = \sum_{i=1}^m \lambda_i C_{ij} + \mu = C_{oi}$$

$$\sum_{i=1}^m \lambda_i = 1$$

که در آن μ ضریب لاگرانژ، C_{ij} و C_{oi} نیز مشابه نمادهای معادله کریگینگ بتا-دوجمله‌ای (فصل دوم) است.

گزینه دوم، برآورد کریگینگ بدون شرط اریبی روی Y است. کریگینگ بدون شرط اریبی وابسته به شرایط زیر است:

$$E[\hat{Y}(s_o) | Y(s_o)] = Y(s_o).$$

لاجونی (۱۹۹۱) نشان داد که

$$E[Y(s_i) | Y(s_o)] = \mu + \frac{C_{oi}}{\sigma_Y^2} \{Y(s_o) - \mu\}$$

بنابراین علاوه بر یک بودن مجموع وزن‌ها، برای اطمینان از ناریبی شرط اضافی

$$\sum_{i=1}^m \lambda_i C_{oi} = \sigma_Y^2$$

را داریم. بنابراین دستگاه کامل معادلات کریگینگ عبارت است از

$$\sum_{i=1}^m \lambda_i C^{\frac{Z_i}{n_i}}(s_i, s_j) + \psi_1 C^Y(s_o, s_i) + \psi_2 = C_{oi}(s_o, s_i)$$

$$\sum_{i=1}^m \lambda_i = 1,$$

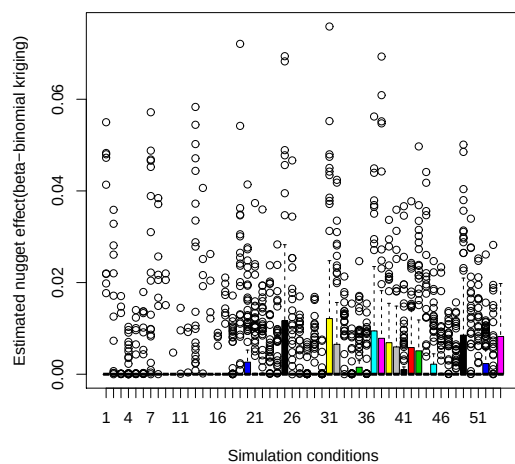
$$\sum_{i=1}^m \lambda_i C^Y(s_o, s_i) = \sigma_Y^2.$$

۲.۱.۴ برآورد پارامترهای هم‌کریگینگ دوجمله‌ای

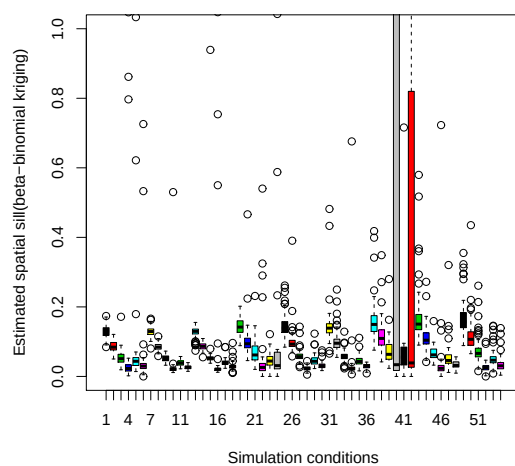
برای ارزیابی عملکرد برآورد پارامترها، برای هر سناریو ۱۰۰ مجموعه داده شبیه‌سازی کردیم. همان‌طور که قبلاً بیان شد، با توجه به نتایج حاصل از شبیه‌سازی، برای برآورد پارامترها از تابع زیان $L_2(\cdot)$ و تعریف تعداد نقاط همسایه $N^*(h)$ ، در استراتژی دوم، استفاده کردیم. برآوردهای پارامترهای تغییرنگار برای هر دو مدل بتا-دوجمله‌ای و دوجمله‌ای از می‌نیم کردن تابع زیان $L_2(\theta)$ ، که در تغییرنگار تجربی و تغییرنگار واقعی (کروی) مقداردهی شده است، به‌دست آمده‌اند. نمودارهای جعبه‌ای برآوردهای پارامترهای اثر قطعه‌ای، ازاره و دامنه برای مدل بتا-دوجمله‌ای در شکل ۱.۴ و برای مدل دوجمله‌ای در شکل ۲.۴ نشان داده شده‌اند. در هر دو شکل، ۱۸ سناریوی اول مربوط به دامنه ۵، ۱۸ سناریو دوم مربوط به دامنه ۱۰ و ۱۸ سناریو سوم مربوط به دامنه ۱۵ هستند. همچنین ۶ سناریو اول هر دامنه فضایی، اندازه نمونه کوچک، ۶ سناریو دوم اندازه نمونه متوسط و ۶ سناریو سوم اندازه نمونه بزرگ را نشان می‌دهند. پارامترهای واقعی دامنه فضایی به ترتیب برای این سناریوها ۵، ۱۰ و ۱۵ می‌باشند. اثر قطعه‌ای واقعی صفر است که تقریباً در تمام شبیه‌سازی‌ها به درستی برآورد شده است. اثر قطعه‌ای مدل دوجمله‌ای حتی کوچک‌تر از اثر قطعه‌ای کریگینگ بتا-دوجمله‌ای برآورد شده است. همچنین با توجه به دو نمودار جعبه‌ای مربوط به برآورد این پارامتر در هر دو مدل، می‌توان گفت اندازه نمونه در برآورد اثر قطعه‌ای تأثیری ندارد. در واقع الگوی قابل تشخیصی در شکل معلوم نیست که این نشان می‌دهد (۱) اصلاح تغییرنگار بتا-دوجمله‌ای، برآورد پارامترهای فرآیند بتای پنهان را بهبود بخشیده است؛ (۲) اصلاحی مشابه در تغییرنگار دوجمله‌ای نیز بهبودی مشابه را در پی داشته است.

نمودارهای جعبه‌ای برآورد پارامتر ازاره در هر دو مدل، چوله به راست هستند و در مدل بتا-دوجمله‌ای تعداد نقاط پرت کم‌تری نسبت به اثر قطعه‌ای وجود دارد. با افزایش دامنه فضایی نقاط پرت بیش‌تر هم می‌شوند. واریانس برآورد این پارامتر در مدل دوجمله‌ای خیلی بزرگ‌تر از مدل بتا-دوجمله‌ای است. مقادیر میانه ازاره فضایی تقریباً همیشه کوچک‌تر از تغییرپذیری میدان زیربنایی بتا است. در حقیقت تفاوت بین میانه ازاره و ازاره فضایی واقعی شدید است. این نشان می‌دهد که هم‌کریگینگ دوجمله‌ای ممکن است با همگرایی در تقابل باشد، زیرا در بسیاری از موارد برآوردهای ازاره برای داده‌ها بسیار غیرمنطقی هستند. در بعضی سناریوها، بیش از ۲۵ درصد شبیه‌سازی‌ها مقدار برآوردشده ازاره کاملاً بی‌معنی دارند. از طرفی، برای مدل بتا-دوجمله‌ای نیز برآوردهای ازاره بزرگ‌تر از یک وجود دارند که با توجه به کراندار بودن توزیع بتا در فاصله صفر و یک طبیعی نیست و نشان می‌دهد همگرایی برای بعضی از توزیع‌ها ممکن است مشکل‌ساز شود.

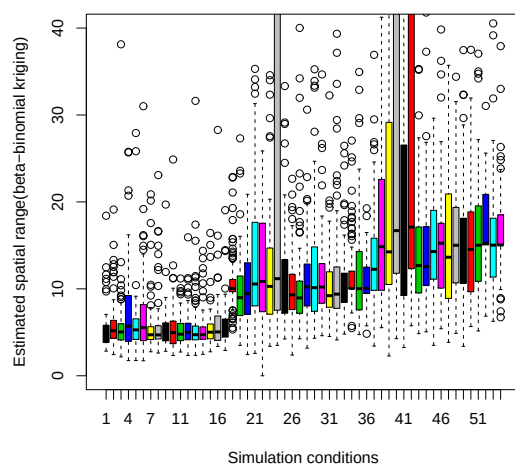
برای پارامتر دامنه فضایی، در مدل بتا-دوجمله‌ای، هر سه مقدار در نظر گرفته‌شده به خوبی بازایی شده‌اند. روندی تقریباً مشابه برای مدل دوجمله‌ای برقرار است اما مشابه ازاره، واریانس برآوردها نسبت به مدل بتا-دوجمله‌ای خیلی بزرگ‌تر است.



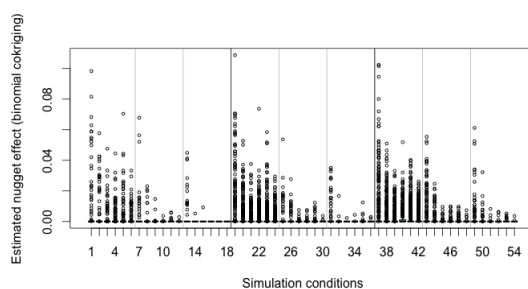
(آ) اثر قطعه‌ای



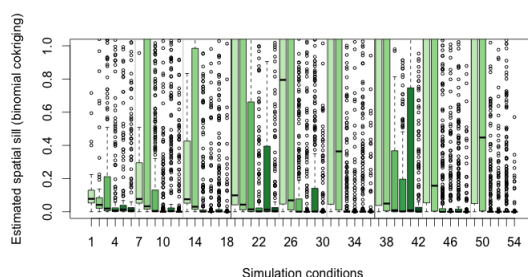
(ب) ازاره فضایی



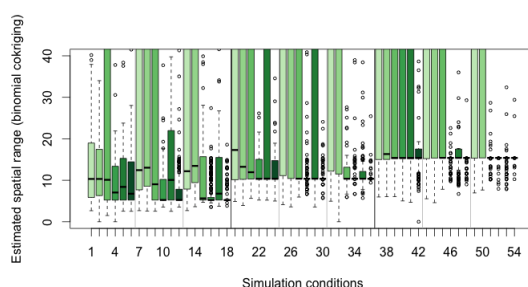
(ج) دامنه فضایی



(آ) اثر قطعه‌ای



(ب) ازاره فضایی



(ج) دامنه فضایی

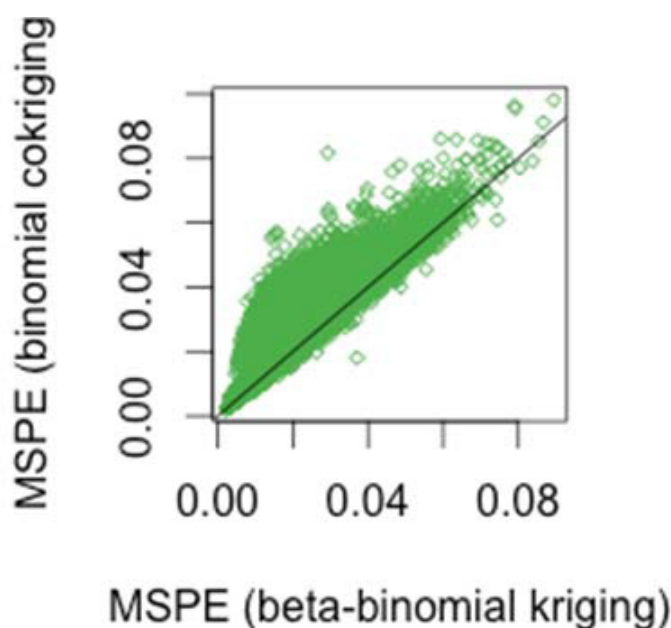
شکل ۲.۴: برآورد پارامترهای تغییرنگار دوجمله‌ای

با توجه به برآوردهای پارامترها در مدل هم‌کریگینگ دوجمله‌ای، می‌توان گفت که نگرانی جدی در مورد استفاده از این مدل برای پیش‌گویی فضایی وجود دارد. به نظر می‌رسد در این مدل، بهترین سناریوی استفاده از آن سناریویی با اندازه نمونه بزرگ، دامنه فضایی نسبتاً بزرگ و توزیع زیربنایی بتا با تغییرپذیری کم است.

۳.۱.۴ پیش‌گویی فضایی

برای بررسی عملکرد پیش‌گویی دو روش معرفی شده، ۱۰۰ مجموعه داده را (به ازای هر سناریوی شبیه‌سازی) بر روی یک شبکه 50×50 شامل ۲۵۰۰ گره تولید کردیم و با استفاده از دو روش کریگینگ بتا-دوجمله‌ای و هم‌کریگینگ دوجمله‌ای (با پارامترهای برآورد شده تغییرنگار)

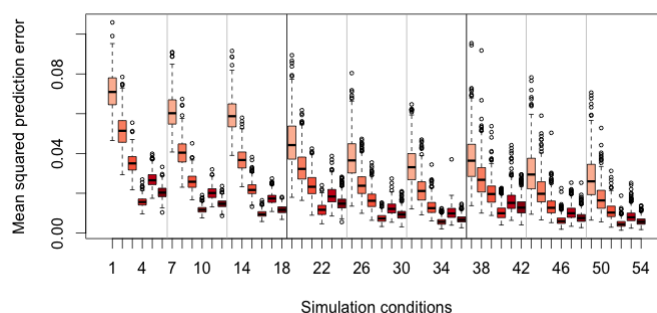
مقادیر نسبت‌های دوجمله‌ای را، برای هر مجموعه داده، پیش‌گویی کردیم. سپس میانگین توان دوم خطای پیش‌گویی^۲ (MSPE) را برای دو مدل بتا-دوجمله‌ای و دوجمله‌ای محاسبه کردیم. شکل ۳.۴ نمودار پراکنش مقادیر MSPE دو روش را نشان می‌دهد. این شکل مقادیر MSPE تجمیع‌شده همه سناریوها را با هم نشان می‌دهد. از این شکل واضح است که مدل کریگینگ بتا-دوجمله‌ای، بر حسب معیار MSPE، عملکرد کاملاً برتری نسبت به مدل هم کریگینگ دوجمله‌ای دارد. در واقع، در ۷۹/۴ درصد شبیه‌سازی‌ها، کریگینگ بتا-دوجمله‌ای MSPE کم‌تری از هم کریگینگ دوجمله‌ای تولید کرده است. در مواردی هم که هم کریگینگ دوجمله‌ای MSPE کوچک‌تری دارد، اختلاف نسبتاً ناچیز است.



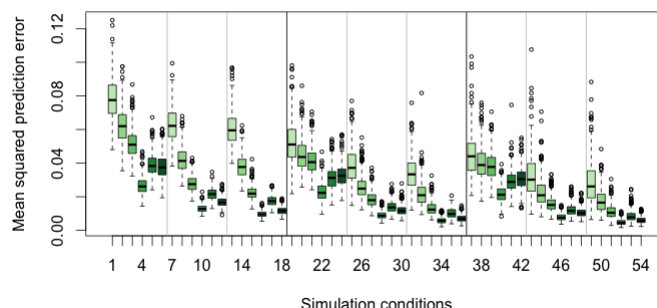
شکل ۳.۴: مقایسه مقادیر MSPE برای دو روش کریگینگ بتا-دوجمله‌ای و هم کریگینگ دوجمله‌ای

برای بررسی تاثیر حجم نمونه و اندازه دامنه فضایی بر قدرت پیش‌گویی، مقادیر MSPE را برای دو مدل به تفکیک هر ۵۴ سناریو در شکل ۴.۴ گزارش کرده‌ایم. با توجه به این دو شکل اختلاف توانایی دو روش، به‌ویژه در اندازه نمونه کوچک، مشهود است. ولی دامنه تاثیر چندانی ندارد. بنابراین می‌توان نتیجه‌گیری کرد که برای اندازه‌های نمونه کوچک و متوسط دقت کریگینگ بتا-دوجمله‌ای بیشتر است و باید از آن استفاده کرد.

^۲ Mean Squared Prediction Errors



(آ) میانگین توان دوم خطای پیش‌گویی مدل بتا- دوجمله‌ای



(ب) میانگین توان دوم خطای پیش‌گویی مدل دوجمله‌ای

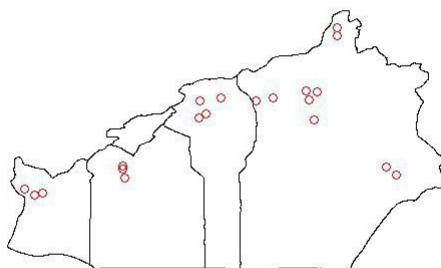
شکل ۴.۴: میانگین توان دوم خطای پیش‌گویی دو مدل به ازای ۵۴ سناریوی شبیه‌سازی

۲.۴ پیش‌گویی و نقشه پهنه‌بندی نرخ بارش در استان سمنان

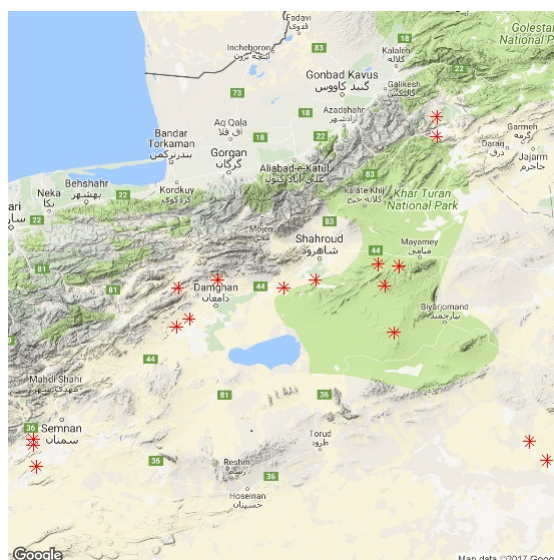
در این بخش کاربرد روش پیش‌گویی معرفی شده را بر روی یک مجموعه داده واقعی مورد بررسی قرار می‌دهیم. مجموعه داده مورد نظر مربوط به تعداد روزهای دارای بارندگی در استان سمنان است که به مدت ۵ ماه از اول آبان تا پایان اسفند سال ۱۳۹۱ طی ۱۴۹ روز در ۲۰ ایستگاه هواشناسی استان سمنان ثبت شده‌اند. با توجه به این که مشاهدات تعداد روزهای دارای بارندگی است و داده‌ها به موقعیت قرار گرفتن شان در محیط مورد مطالعه بستگی دارند، بنابراین پاسخ‌ها به صورت گسسته و شمارشی هستند و ماهیت فضایی دارند. هدف از تحلیل این داده‌ها، تهیه نقشه پهنه‌بندی پیش‌گویی نرخ بارش در کل سطح استان با استفاده از مدل منعطف بتا- دوجمله‌ای فضایی است. موقعیت ایستگاه‌های هواشناسی بر روی نقشه استان

سمنان در شکل‌های ۵.۴ و ۶.۴ مشخص شده‌اند. با توجه به بارش بیش‌تر در شمال استان، تعداد ایستگاه‌های هواشناسی در کرانه شمالی استان سمنان ۱۸ و در کرانه جنوب شرقی آن فقط ۲ ایستگاه است.

Location of Stations



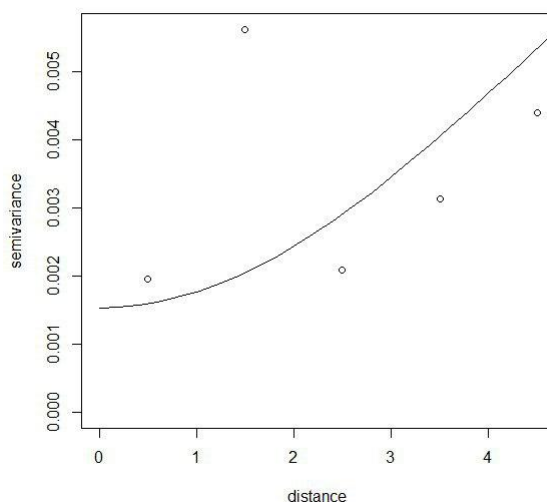
شکل ۵.۴: موقعیت ایستگاه‌های هواشناسی استان سمنان



شکل ۶.۴: موقعیت ایستگاه‌های هواشناسی استان سمنان

برای برآورد ساختار همبستگی روزهای دارای بارش در استان سمنان طی ۱۴۹ روز ثبت‌شده در سال ۱۳۹۱، ابتدا تغییرنگار داده‌های بارندگی را برآورد کردیم. برای استفاده از یک مدل همبستگی پارامتری، مدل‌های مختلف همبستگی فضایی را به تغییرنگار برآورد شده

برازش دادیم. بهترین مدل برازش شده که یک مدل گاوسی است در شکل ۷.۴ نشان داده شده است.

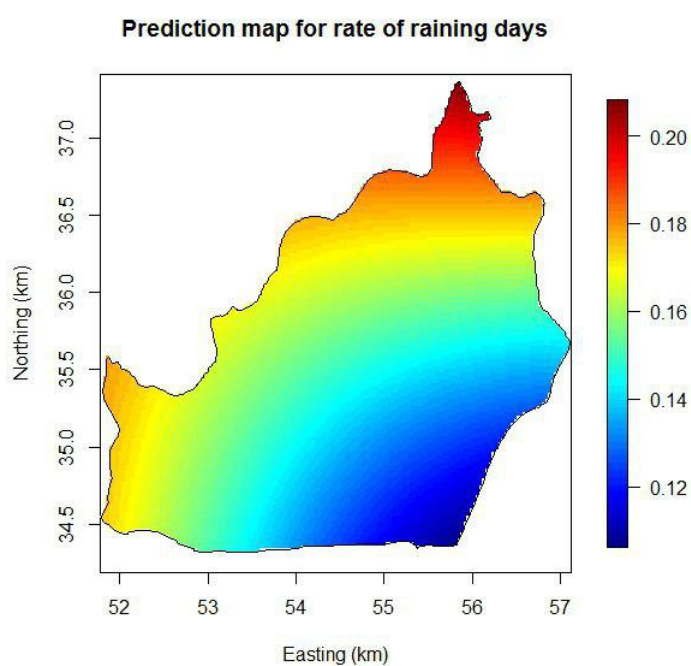


شکل ۷.۴: نیم‌تغییرنگار تجربی و برازش مدل گاوسی به داده‌ها

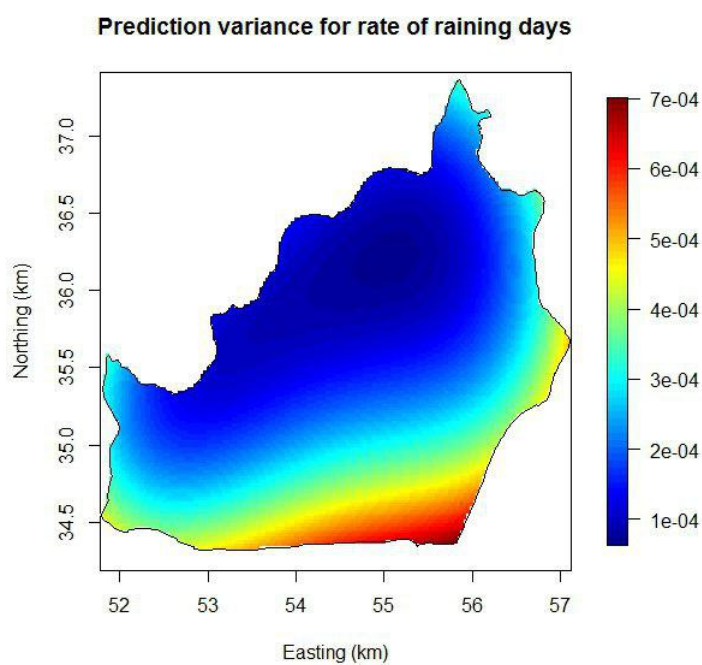
برای پیش‌گویی نقاط فاقد مشاهده در استان از یک شبکه $5^\circ \times 5^\circ$ که کران‌های استان را کاملاً پوشش می‌دهند، استفاده کردیم. تابع تغییرنگار فضایی را نیز مدل گاوسی برازش‌شده بر تابع تغییرنگار تجربی در نظر گرفتیم. متغیر پاسخ، تعداد روزهای دارای بارندگی است که فرض کردیم دارای توزیع دوجمله‌ای شرطی مستقل باشد و متغیر پنهان آن نیز دارای توزیع بتا است. متأسفانه این مجموعه داده شامل متغیرهای رگرسیونی نبود که از اطلاعات آن‌ها برای پیش‌گویی کاراتر بتوان استفاده کرد. شکل ۸.۴ نقشه پهنه‌بندی پیش‌گویی نرخ روزهای دارای بارندگی را روی کل استان سمنان بر اساس مدل بتا-دوجمله‌ای نشان می‌دهد. همان‌طور که از روی شکل واضح است، هر چه به سمت شمال استان حرکت می‌کنیم، نرخ بارندگی نیز افزایش می‌یابد به‌طوری که این نرخ به بیش از 20° درصد نیز می‌رسد. یکی از دلایل این ویژگی همجواری با استان‌های شمالی کشور یعنی گلستان و مازندران است.

شکل ۹.۴ نیز واریانس نرخ بارش پیش‌گویی را نمایش می‌دهد با توجه به آن که اکثر اطلاعات در کرانه شمالی استان ثبت شده‌اند، واریانس پیش‌گویی نیز در کرانه شمالی کوچک‌تر از کرانه جنوبی و مرکز استان می‌باشد.

از نقشه پهنه‌بندی پیش‌گویی نرخ بارش، می‌توان برای تصمیم‌گیری‌های مدیریت آب استفاده کرد. این نقشه، وضعیت بارش در سطح استان را برای تمام نقاط فاقد ایستگاه هواشناسی، به‌صورت تقریبی، نمایش می‌دهد.



شکل ۸.۴: نقشه پهنه‌بندی پیش‌گویی نرخ بارش در استان سمنان



شکل ۹.۴: واریانس نرخ پیش‌گویی بارش در استان سمنان

فصل ۵

تحلیل فضایی داده‌های شمارشی مشبکه‌ای با مدل اتوبتا-دوجمله‌ای

در این فصل به تعمیم مدل بتا-دوجمله‌ای برای تحلیل فضایی داده‌های مشبکه‌ای در چارچوب استنباط بیزی می‌پردازیم. با توجه به پیچیدگی توزیع پسین این مدل، از یک روش تقریبی بیزی به نام تقریب لاپلاس آشیانی جمع‌بسته^۱ (INLA) استفاده می‌کنیم. این روش توسط رو و همکاران (۲۰۰۹) برای استنباط در مدل‌های گاوسی پنهان^۲ معرفی شد که مدل‌های فضایی را نیز شامل می‌شود. استفاده از روش INLA در برازش مدل‌های فضایی و فضایی-زمانی رو به افزایش است که می‌توان به منابع اشاره‌شده در میرزاجانی (۱۳۹۵) مراجعه کرد. در ادامه داده‌های مربوط به بیماری سرطان معده در استان گیلان را با استفاده از روش INLA تحلیل می‌کنیم. هدف از تحلیل این داده‌ها، تهیه نقشه پهنه‌بندی نرخ سرطان معده در استان گیلان با استفاده از مدل منعطف بتا-دوجمله‌ای و همچنین مقایسه آن با مدل دوجمله‌ای است.

^۱ Integrated nested Laplace approximation

^۲ Latent Gaussian models

۱.۵ مقدمه

در بسیاری از مطالعات، تحلیل داده‌های فضایی که به صورت شمارشی یا متوسط‌گیری روی نواحی دلخواه فضای مورد مطالعه هستند، مورد نظر است. این داده‌ها، داده‌های شبکه‌ای نامیده می‌شوند. برای مدل‌بندی این داده‌ها از یک میدان تصادفی مانند

$$Z(\cdot) = \{Z(s) : s \in D\}$$

استفاده می‌شود که در آن D زیرمجموعه‌ای ثابت و شمارش‌پذیر (معمولاً متناهی) از $\mathbb{R}^d$ می‌باشد، همچنین D ممکن است به طور منظم یا نامنظم شبکه‌بندی شده باشد. تصاویر ماهواره‌ای از سطح زمین که در آن‌ها سطح زمین به تعداد مساوی مستطیل کوچک به صورت شبکه منظم در $\mathbb{R}^2$ هستند، مثالی برای داده‌های شبکه‌ای منظم است و تعداد افراد مبتلا به سرطان در تمام مناطق خدمات درمانی کشور یک مثال برای داده‌های شبکه‌ای نامنظم است.

در حقیقت میدان تصادفی شبکه‌ای را می‌توان به صورت $\{Z(s_i) : s_i \in \{s_1, \dots, s_n\}\}$ نمایش داد که در آن $s_1, \dots, s_n$ یک افراز از ناحیه D است، یعنی

$$\bigcup_{i=1}^n s_i = D \quad s_i \cap s_j = \emptyset \quad i \neq j.$$

برای درک بهتر، اگر D مساحت استان گیلان و $s_1, \dots, s_{16}$ شهرستان‌های گیلان باشند، در این صورت $Z(s_1), \dots, Z(s_{16})$ را می‌توان تعداد افراد مبتلا به سرطان معده، تعداد حوادث رانندگی یا نرخ رشد اقتصادی هر شهرستان در نظر گرفت. این نوع داده‌ها روش تحلیل مخصوص خود را دارند. در تحلیل داده‌های شبکه‌ای، هدف اصلی مدل‌بندی مشاهدات است، در صورتی که در زمین‌آمار پیش‌گویی متغیر در یک موقعیت جدید فاقد مشاهده هدف اصلی محسوب می‌شود. از آنجایی که داده‌های شبکه‌ای در یک زیرمجموعه شمارش‌پذیر از $\mathbb{R}^d$ امکان تحقق و مشاهده دارند، بنابراین باید دقت شود برای نقاطی می‌توان پیش‌گویی کرد که امکان تحقق دارند و در غیر این صورت نتیجه بی‌معنی خواهد بود. در متون مختلف علمی از عبارت‌هایی مانند داده‌های ناحیه‌ای و منطقه‌ای^۳ نیز برای داده‌های شبکه‌ای استفاده می‌شوند.

برای این که ساختار فضایی داده‌های شبکه‌ای را مدل‌بندی کنیم، دو رهیافت متفاوت وجود دارند که تفاوت این دو رهیافت به مجموعه اندیس‌گذار D در میدان تصادفی

$$\{Z(s) : s \in D\}$$

برمی‌گردد. در رهیافت اول، فرض می‌شود که مشاهدات از یک مجموعه اندیس‌گذار پیوسته آمده‌اند. همچنین فرض می‌شود داده‌ها برای هر ناحیه در مرکز ناحیه مشاهده شده‌اند. در

^۳Regional data

واقع خلاصه اطلاعات تمام ناحیه در مرکز ناحیه است، یعنی اطلاعات کل ناحیه به یک نقطه منتسب می‌شود. در این رهیافت مدل‌بندی داده‌ها با اندیس پیوسته صورت می‌گیرد، در صورتی که ماهیت داده‌ها شبکه‌ای است و در هر نقطه امکان تحقق وجود ندارد. در حقیقت در این رهیافت، روش زمین‌آمار که مختص داده‌های با اندیس پیوسته است برای داده‌های با اندیس گسسته به کار برده می‌شود. دلیل استفاده از این رهیافت، سر راست بودن تابع کوواریانس فضایی و معمولاً تفسیر راحت آن می‌باشد (ملانی، ۲۰۰۲).

در رهیافت دوم، برای مدل‌بندی ساختار فضایی داده‌های شبکه‌ای، طبیعت گسسته مکان داده‌های شبکه‌ای در نظر گرفته می‌شود و ساختار همسایگی مشاهدات بر اساس شکل شبکه‌ها تعریف می‌شود. بنابراین به جای اندازه‌گیری فاصله بین مراکز ناحیه‌ها یک روش دیگر بر اساس همسایگی ناحیه‌ها تعریف می‌شود. به عنوان مثال از میزان مرز مشترک ناحیه‌ها در ساختار استفاده می‌کنند. در حقیقت وقتی ساختار همسایگی به این صورت تعریف می‌شود، مدل‌هایی مانند اتورگرسیو در سری‌های زمانی تعریف می‌شود. با این دیدگاه ویتل (۱۹۵۴) مدل اتورگرسیو همزمان، هاینینگ (۱۹۷۸) مدل میانگین متحرک، بارتلت (۱۹۷۱) و بیسگ (۱۹۷۴) مدل‌های اتورگرسیو شرطی^۴ (CAR) را برای تحلیل داده‌های شبکه‌ای ارائه کردند.

۲.۵ داده‌های فضایی شبکه‌ای

بیرخوف و هامرسل (۱۹۶۷) و مازارینو (۱۹۸۳) عبارت شبکه را اولین بار در متون ریاضی مطرح کردند، سپس بیسگ (۱۹۷۵، ۱۹۷۴) از آن در تحلیل داده‌های شبکه‌بندی شده استفاده کرد. بیسگ (۱۹۷۴) مدل‌های اتوگوسی^۵، اتودوجمله‌ای^۶، اتولجستیک^۷، اتوپواسون^۸ را برای تحلیل داده‌های شبکه‌ای در حالت یک متغیره به کار برد، که در همه آن‌ها از میدان‌های تصادفی گاوسی استفاده کرده است. جاکسون (۲۰۰۳) روش‌های مختلف برآورد پارامترها و روش‌های مختلف شبیه‌سازی این مدل‌ها را مورد مطالعه قرار داده و سپس آن‌ها را مقایسه کرده است. چن (۲۰۰۲) تحلیل بیزی مدل اتولجستیک را برای تحلیل توزیع گونه خاصی از گیاهان مطرح کرده است. از بین این مدل‌ها مدل اتوگوسی یا اتورگرسیو شرطی گاوسی به دلیل وجود صورت بسته درست‌نمایی از لحاظ کاربردی مورد توجه قرار گرفته است. حالت چندمتغیره آن اولین بار توسط ماردیا (۱۹۸۸) برای پردازش تصاویر معرفی گردید و سپس بیلهمیر و همکاران (۱۹۹۷)، پتیت و همکاران (۲۰۰۲)، کارلین و بنرجی (۲۰۰۳)، و گلفند و وناتسو (۲۰۰۳) در مورد آن مطالعات بیش‌تری انجام دادند. تلاش آن‌ها استفاده از رهیافت بیزی به‌ویژه بیزی سلسله‌مراتبی برای برآورد پارامترهای مدل اتوگوسی چندمتغیره بوده

^۴ Conditional autoregressive

^۵ Auto Gaussian

^۶ Auto binomial

^۷ Auto logistic

^۸ Aoto Poisson

است. ساین و کرسی (۲۰۰۴) مدل اتوگوسی چندمتغیره را در حالتی که ماتریس همبستگی متقابل فضایی نامتقارن است، ارایه کرده‌اند. اغلب در متون از عنوان میدان تصادفی مارکوفی چندمتغیره^۹ به جای مدل اتوگوسی چندمتغیره استفاده شده است.

همان گونه که در سری‌های زمانی رابطه هر مقدار جدید در آینده با مشاهدات گذشته نزدیک مدل‌بندی و مقدار آن پیش‌گویی می‌شود، در میدان تصادفی فضایی شبکه‌ای نیز هر موقعیت جدید بر اساس داده‌های همبسته با آن، یعنی مشاهدات واقع در نواحی همسایه نزدیک، مدل‌بندی و در بعضی موارد پیش‌گویی می‌شود. بنابراین به منظور مدل‌بندی داده‌های ناحیه‌ای $y_i, i = 1, \dots, n$ ، لازم است اثر سایر همسایه‌ها بر این ناحیه مشخص شود. مجموعه همسایگی موقعیت i با N_i نمایش داده می‌شود و عبارت است از مجموعه موقعیت‌هایی که بیش‌ترین اطلاع را در خصوص مشاهده y_i در اختیار دارند. همسایگی یک ناحیه بر این اساس که منظم یا نامنظم باشد به صورت دو موقعیت هم‌مرز یا دو ناحیه که در یک فاصله معین از هم قرار دارند، تعریف می‌شود. اغلب، همسایگی یک مکان به صورت تجربی و بر اساس فاصله تعیین می‌شود. مثال‌های زیر روش‌های مختلف تعیین همسایگی را نشان می‌دهند.

- در مثال معروف داده‌های نشانگان مرگ و میر ناگهانی نوزادان^{۱۰} (SIDS) در کارولینای شمالی، نواحی یا چندضلعی‌ها توسط طول و عرض جغرافیایی مشخص شده‌اند و نواحی که در شعاع ۳۰ مایلی ناحیه i قرار دارند، همسایه آن تلقی می‌شوند. منطقه کارولینا مثالی از شبکه نامنظم با پوشش چند ضلعی است.
- برای بررسی میزان کود شیمیایی در یک زمین کشاورزی که به کرت‌های هم‌اندازه و منظم تقسیم شده است، همسایگی هر خانه ممکن است توسط همسایه‌ها از مرتبه اول (کرت‌های هم‌جوار با هر کرت) یا از مرتبه اول و دوم در نظر گرفته باشند. در جدول ۱.۵ برای مکان $\times$ ، همسایه‌های مرتبه اول با علامت $\bullet$ و همسایه‌های مرتبه دوم با علامت $\otimes$ مشخص شده‌اند.

			$\otimes$			
			$\bullet$			
	$\otimes$	$\bullet$	$\times$	$\bullet$	$\otimes$	
			$\bullet$			
			$\otimes$			

جدول ۱.۵: شبکه منظم با پوشش شبکه‌ای از مربع‌ها

^۹ Markov random fields

^{۱۰} Sudden infant death syndrome

- در مدل‌های اقتصادی، همسایگی ممکن است بر اساس تبادلات اقتصادی بین دو ناحیه مشخص شود. به عنوان مثال اگر متغیر مورد مطالعه، تولیدات اقتصادی کشورهای مختلف باشد، همسایه ناحیه i ام، کشورهایی هستند که با آن مبادلات مثبت تجاری دارند. بنابراین دو کشور همسایه، لزوماً مرز مشترک جغرافیایی ندارند.

۱.۲.۵ روش‌های انتساب وزن‌های همسایگی

پس از آن که نوع همسایگی را مشخص کردیم، بیان تاثیر این همسایگی دارای اهمیت است. این تاثیر از طریق انتساب وزن به هر یک از همسایه‌ها به دست می‌آید. برای نشان دادن وزن‌های همسایگی از ماتریس مجاورت^{۱۱} C استفاده می‌کنیم که در آن درایه c_{ij} میزان تاثیر ناحیه i روی ناحیه j را نمایش می‌دهد. همچنین $c_{ii} = 0$ و $c_{ij} = c_{ji}$. ماتریس C را می‌توان به صورت

$$C = \left(\frac{c_{ij}}{c_{i+}} \right)$$

به یک ماتریس استاندارد تبدیل کرد که در آن $c_{i+} = \sum_j c_{ij}$. برای مشخص کردن ماتریس C ، می‌توان از ویژگی‌های مختلفی مانند فاصله بین مراکز نواحی، مرز مشترک نواحی، جمعیت هر ناحیه، مساحت نواحی و میزان مبادلات فرهنگی و تجاری نواحی استفاده کرد. ساده‌ترین حالت، انتخاب ماتریس مجاورت به صورت

$$c_{ij} = \begin{cases} 1 & j \in N_i \\ 0 & i = j \end{cases}$$

است. یکی از راه‌های انتساب وزن‌ها، استفاده از فاصله بین نواحی است. به عبارت دیگر، تعریف تابع نزولی بر اساس فاصله به گونه‌ای که به همسایه‌های نزدیک، وزن بزرگ‌تر و به همسایه‌های دور، وزن کوچک‌تری اختصاص می‌دهد. مساحت‌های نواحی نیز می‌توانند به این صورت برای تعریف درایه‌های ماتریس C استفاده شوند به طوری که $c_{ij} = \frac{A_{ij}}{\sum_{j \in N_i} A_{ij}}$ ، که در آن A_{ij} مساحت j امین همسایه ناحیه i است. در صورتی که جمعیت نواحی معلوم باشد، می‌توان وزن‌ها را به صورت

$$c_{ij} = \frac{A_{ij}}{\sum_{j \in N_i} A_{ij}} \times \frac{n_j}{n_i}$$

تعدیل کرد، به طوری که n_i و n_j به ترتیب جمعیت ناحیه i ام و جمعیت ناحیه j ام هستند. همچنین با در نظر گرفتن فاصله بین نواحی، می‌توان وزن‌های c_{ij} را به صورت

$$c_{ij} = \frac{A_{ij}}{\sum_{j \in N_i} A_{ij}} \times \frac{n_j}{n_i} \times d'_{ij}$$

^{۱۱}Adjacency matrix

نوشت که در آن d_{ij} فاصله بین مرکز دو ناحیه و γ پارامتری است که می‌بایست برآورد شود. یک روش دیگر برای تعریف ماتریس C استفاده از مرز مشترک نواحی است، به عبارت دیگر، می‌توان نوشت

$$c_{ij} = \left(\frac{L_{ij}}{L_i} \right)^\tau \quad (1.5)$$

که در آن L_{ij} طول مرز مشترک ناحیه i با ناحیه j ، L_i محیط ناحیه i و $\tau \geq 0$ پارامتری است که باید برآورد شود. حالت کلی‌تری از روش (۱.۵) در نظر گرفتن درایه‌ها به صورت

$$c_{ij} = \frac{L_{ij}}{L_i} d_{ij}^{-\gamma\tau} \quad (2.5)$$

است، که در آن هم از فاصله بین نواحی و هم از طول مرزهای مشترک دو ناحیه استفاده می‌شود.

کرسی (۱۹۹۳) ضمن معرفی تابع همسایگی (۲.۵)، تعریف وزن‌های همسایگی را بر اساس تابع مبتنی بر فاصله که در آن به همسایه‌های دور، وزن کوچک‌تر و به همسایه‌های نزدیک وزن بزرگ‌تری، به صورت

$$C_{ij} = \begin{cases} \rho \frac{d_{ij}^{-k} n_i}{c(k) n_j} & j \in N_i \\ 0 & j \notin N_i \end{cases}$$

داده می‌شود، پیشنهاد کرد که در آن ρ مقداری ثابت، d_{ij} فاصله بین ناحیه i و همسایه j ام، n_i و n_j جمعیت‌های دو ناحیه i و j ، N_i مجموعه همسایگی‌های ناحیه i و

$$c(k) = \max\{d_{ij}^{-k}, j \in N_i, i = 1, 2, \dots, n\}$$

عامل مقیاس است که برای بدون واحد کردن وزن‌های c_{ij} به کار می‌رود. همچنین k مقداری ثابت است که معمولاً مقادیر ۰، ۱ یا ۲ را برای آن در نظر می‌گیرند.

۳.۵ مدل‌بندی داده‌های شبکه‌ای

برای مدل‌بندی داده‌های شبکه‌ای معمولاً به دو نوع تغییرات توجه می‌شود. تغییرات بزرگ مقیاس که در آن میانگین به دلیل وابستگی موقعیت‌های فضایی به سایر موقعیت‌های تبیینی تغییر می‌کند و تغییرات کوچک مقیاس که به دلیل اثرات متقابل همسایگی‌ها ایجاد می‌شود. معمولاً شکل کلی

$$Z(s) = \mu(s) + \epsilon, \quad s \in D$$

برای مدل‌بندی داده‌های شبکه‌ای تعریف می‌شود، که به آن مدل رگرسیون فضایی نیز می‌گویند و در آن $Z(s)$ مقدار فرآیند تصادفی در موقعیت s است، $\mu(s)$ میانگین میدان تصادفی در موقعیت s است، که ممکن است ثابت یا تابعی از متغیرهای تبیینی باشد و ϵ بیانگر تغییرات

کوچک مقیاس است که معمولاً توزیع نرمال $\epsilon \sim N(0, \Sigma)$ برای آن در نظر گرفته می‌شود و Σ با ابعاد $n \times n$ ماتریس واریانس متغیر تصادفی در همه موقعیت‌هاست. در واقع تغییرات کوچک مقیاس در Σ منعکس می‌گردد. انتخاب‌های مختلف برای ساختار Σ موجب روش‌های مختلف مدل‌بندی داده‌های شبکه‌ای به صورت اتورگرسیو شرطی CAR، اتورگرسیو همزمان^{۱۲} (SAR) و میانگین متحرک MA شده است. در همه این مدل‌ها، پذیره توزیع نرمال چندمتغیره مورد نیاز است و ساختار کوواریانس Σ در هر کدام از آن‌ها به صورت زیر در نظر گرفته می‌شود:

$$CAR: \Sigma = (I - \rho W)^{-1} D \sigma^2$$

$$SAR: \Sigma = [(I - \rho W)' D (I - \rho W)]^{-1} \sigma^2$$

$$MA: \Sigma = (I - \rho W) D (I - \rho W)' \sigma^2$$

که در آن‌ها ρ و σ^2 پارامترهای تغییرات کوچک مقیاس‌اند، W ماتریس وزن‌های همسایگی و D ماتریس قطری برای منظور کردن واریانس ناهمگن توزیع‌های کناری است. ماتریس D گاهی ماتریس وزن نامیده می‌شود، که در این صورت نباید با ماتریس W اشتباه شود. دو مدل CAR و SAR مشابه فرآیند اتورگرسیو در سری‌های زمانی تعریف می‌شوند. در مدل‌های مذکور اگر ماتریس Σ متقارن باشد، سبب سادگی بسیاری از محاسبات می‌شود. معمولاً از ماتریس D برای متقارن کردن Σ استفاده می‌شود. به عنوان مثال اگر

$$Var(Z(s_i) | \{Z(s_j); j \in N_i\}) = \sigma_i^2 = \frac{\sigma^2}{n_i}$$

باشد، در این صورت با انتخاب ماتریس قطری D با عناصر قطری $\frac{1}{n_i}$ ماتریس Σ متقارن می‌شود.

۴.۵ معرفی مدل‌های اتورگرسیو شرطی

یکی از معروف‌ترین مدل‌ها برای تحلیل داده‌های شبکه‌ای، مدل CAR است که به وسیله توزیع‌های احتمال شرطی ساخته می‌شود. موقعیت‌های قرارگیری داده‌ها در مدل‌بندی توزیع احتمال نقش اساسی ایفا می‌کنند. بنابراین لازم است نحوه تاثیر این موقعیت‌ها در مدل‌بندی مورد بررسی قرار گیرد. برای این منظور ابتدا ویژگی‌های زنجیر مارکوف را مورد بررسی قرار می‌دهیم. سپس مدل CAR را معرفی می‌کنیم.

۱.۴.۵ زنجیرهای مارکوف

فرض کنید میدان تصادفی $\{Z(s); s = 0, 1, \dots\}$ در طول زمان به دست آمده است و بدون کاستن از کلیت مطلب $Z(\cdot)$ را گسسته در نظر بگیرید. در این صورت زنجیر مارکوف در مدل

^{۱۲} Simultaneously autoregressive

احتمالی زیر صدق می‌کند (کرسی، ۱۹۹۳):

$$P(Z(1), \dots, Z(n) | Z(0)) = \prod_{s=1}^n Q(Z(s), Z(s-1)). \quad (3.5)$$

یعنی احتمال توام متغیرها را می‌توان به صورت حاصل ضرب توابعی وابسته به متغیرهای متوالی نوشت، که به آن ویژگی فاقد حافظه نیز می‌گویند. با توجه به تعریف (۳.۵)

$$\begin{aligned} P(Z(i) | Z(0), \dots, Z(i-1)) &= \frac{P(Z(1), \dots, Z(i) | Z(0))}{P(Z(1), \dots, Z(i-1) | Z(0))} \\ &= \frac{\prod_{s=1}^i Q_s((Z(s), Z(s-1)))}{\sum_{Z(i)} \prod_{s=1}^i Q_s((Z(s), Z(s-1)))} \\ &= \frac{Q_i((Z(i), Z(i-1)))}{\sum_{Z(i)} Q_s((Z(s), Z(s-1)))}. \end{aligned}$$

بنابراین هر مشاهده در زمان i فقط به مشاهده قبلی بستگی دارد. در مدل‌بندی داده‌های شبکه‌ای نیز ویژگی فراموشی زنجیر مارکوف صدق می‌کند، زیرا مقدار مشاهده مکان s_i به طور منطقی فقط به موقعیت‌هایی که در همسایگی این مکان قرار دارند، بستگی دارد.

۲.۴.۵ معرفی مدل

بیسگ (۱۹۷۴) مدل اتورگرسیو شرطی را برای مدل‌بندی داده‌های شبکه‌ای به صورت زیر تعریف کرد:

اگر میدان $\{Z(s_i); i = 1, 2, \dots, n\}$ گاوسی باشد، توزیع شرطی $Z(s_i)$ نسبت به بقیه مشاهدات $Z_{-i} = \{Z(s_j); j \neq i; j \in N_i\}$ به صورت

$$P(Z(s_i) | Z_{-i}) = \frac{1}{\sqrt{2\pi s_i^2}} \exp\left\{-\frac{1}{2} \frac{(Z(s_i) - \theta_i\{Z_{-i}\})^2}{s_i^2}\right\} \quad i = 1, \dots, n \quad (4.5)$$

است که در آن $E(Z(s_i) | Z_{-i}) = \theta_i\{Z_{-i}\}$ و $Var(Z(s_i) | Z_{-i}) = s_i^2$ به ترتیب میانگین و واریانس شرطی هستند. تحت شرایط نظم، میانگین شرطی به صورت

$$\theta_i\{Z_{-i}\} = E(Z(s_i) | Z_{-i}) = \mu_i + \sum_{j \in N_i} c_{ij} (Z(s_j) - \mu_j) \quad (5.5)$$

خواهد بود، که در آن $c_{ii} = 0$ و برای $j \notin N_i$ ، $c_{ij} = 0$. با توجه به (۴.۵) و (۵.۵) و قضیه بیسگ، ثابت می‌شود اگر $(I - C)$ وارون پذیر باشد، آن گاه $Z = (Z(s_1), \dots, Z(s_n))$ دارای توزیع نرمال با بردار میانگین μ و ماتریس واریانس $(I - C)^{-1}M$ است، یعنی

$$Z \sim \text{Gau}(\mu, (I - C)^{-1}M) \quad (6.5)$$

که در آن $\mu = (\mu_1, \dots, \mu_n)$ و $Var(Z(s_i) | Z_{-i}) = s_i^2$ ، $M = \text{diag}(s_1^2, \dots, s_n^2)$ برای خوش تعریف بودن (۶.۵) باید $(I - C)^{-1}M$ متقارن و همیشه مثبت باشد، که شرط $c_{ij}s_j^2 = c_{ji}s_i^2$ تقارن $(I - C)^{-1}M$ را تضمین می‌کند.

در صورتی که در (۵.۵) متغیرهای تبیینی x حضور داشته باشند، تغییرات بزرگ مقیاس را به صورت $\mu = x'\beta$ پارامتری می‌کنیم، که در آن بردار q بعدی از متغیرهای تبیینی و β ضرایب رگرسیونی است. تغییرات فضایی کوچک مقیاس که ماتریس C آن‌ها را در برمی‌گیرد، می‌تواند به صورت $C = \phi G$ نوشته شود، که در آن ماتریس معلوم G ساختار همسایگی را تعریف می‌کند و ϕ شدت وابستگی فضایی را کنترل می‌کند. پارامتر ϕ^2 می‌تواند به عنوان توان دوم همبستگی شرطی فضایی بین همسایگی $Z(s_i)$ و $Z(s_j)$ تفسیر گردد، به انتخاب G و M بستگی دارد (کرسی، ۱۹۹۳).

با فرض هم‌واریانسی شرطی، برای $i = 1, \dots, n$ و $Var(Z(s_i)|Z_{-i}) = s^2$ و $M = s^2 I$ با ثابت نگه داشتن $(I - C)$ ، برآورد ماکسیمم درست‌نمایی β و s^2 به صورت

$$\hat{\beta} = [X'(I - C)X]^{-1} X'(I - C)Z$$

$$\hat{s}^2 = \frac{(Z - X\hat{\beta})' [(I - C)(Z - X\hat{\beta})]}{n}$$

حاصل می‌شوند. اگر برآوردهای فوق در تابع درست‌نمایی جایگزین شوند، می‌توان با ماکسیمم کردن تابع درست‌نمایی نسبت به c_{ij} ها برآوردهای متناظر را به دست آورد.

در مدل CAR، لگاریتم دترمینان $(I - C)$ که در تابع درست‌نمایی ظاهر می‌گردد، حائز اهمیت است. این موضوع در مدل‌های شرطی سری زمانی از اهمیت کم‌تری برخوردار است، زیرا در سری‌های زمانی اغلب n بزرگ است و لگاریتم دترمینان ماتریس کوواریانس به صفر میل می‌کند و از آن صرف نظر می‌شود. در مدل‌های فضایی شبکه‌ای چنین فرضی معمولاً وجود ندارد، به همین دلیل سعی زیادی در حل مساله توسط مکلود (۱۹۷۷)، دنت و مین (۱۹۷۸) و انسلی (۱۹۷۹) صورت پذیرفته است. موران و بیسگ (۱۹۷۵) با در نظر گرفتن حالت خاص مدل‌های شرطی نرمال به محاسبه دترمینان $(I - C)$ پرداخته‌اند.

۳.۴.۵ بررسی بیش‌تر مدل‌های شرطی

در این بخش مدل‌های شرطی و کاربرد آن‌ها در تحلیل داده‌های شبکه‌ای را مورد بررسی قرار می‌دهیم. ابتدا قضیه بسیار مهم تجزیه بیسگ (۱۹۷۴) را بیان می‌کنیم، که با استفاده از آن توزیع توام Z در مدل CAR قابل دستیابی است.

قضیه ۱.۴.۵. (تجزیه بیسگ، ۱۹۷۱) اگر $P(\cdot)$ توزیع توام میدان $\{Z(s_i); i = 1, 2, \dots, n\}$ باشد که در شرط مثبت بودن صدق می‌کند، آن‌گاه

$$\frac{P(Z)}{P(Y)} = \prod_{i=1}^n \frac{P(Y(s_i)|Z(s_1), \dots, Z(s_{i-1}), Y(s_{i+1}), \dots, Y(s_n))}{P(Y(s_i)|Z(s_1), \dots, Z(s_{i-1}), Y(s_{i+1}), \dots, Y(s_n))}$$

که در آن $Z = (Z(s_1), \dots, Z(s_n))$ و $Y = (Y(s_1), \dots, Y(s_n))$ دو تحقق ممکن از میدان Z هستند.

تعریف ۱.۴.۵. (جرگه^{۱۳}). مجموعه‌ای از تمام مکان‌های تک عضوی که هیچ همسایه‌ای ندارند یا مکان‌هایی که تمام اعضای آن با هم همسایه‌اند، جرگه نامیده می‌شود. با وجود n موقعیت و بسته به تعریف ساختار همبستگی، اندازه جرگه‌ها ممکن است ۱، ۲ تا n باشد. اگر جرگه‌ها فقط از مرتبه یک باشند، در این صورت موقعیت‌ها باید از همدیگر مستقل باشند. در مباحث کاربردی جرگه‌های حداکثر تا مرتبه ۲ در نظر گرفته می‌شوند.

برای مدل‌بندی داده‌های فضایی بر اساس میدان تصادفی مارکوفی با فرض این که رویداد صفر در هر موقعیت شانس وقوع داشته باشد، بیسگ (۱۹۷۴) تابع $Q(\cdot)$ را برای داده‌های گسسته به صورت

$$Q(Z) = \ln \left(\frac{P(Z)}{P(\circ)} \right)$$

تعریف کرد، که در آن توزیع توام Z به صورت

$$P(Z) = P(\circ) \exp\{Q(Z)\}$$

قابل بیان است. چون

$$1 = \sum P(Z) = P(\circ) \sum \exp\{Q(Z)\}$$

در نتیجه

$$P(\circ) = \frac{\sum P(Z)}{\sum \exp\{Q(Z)\}} = \frac{1}{\sum \exp\{Q(Z)\}}.$$

در نتیجه توزیع توام Z به صورت

$$P(Z) = \frac{\exp\{Q(Z)\}}{\sum \exp\{Q(Z)\}} \quad (۷.۵)$$

به دست می‌آید. در حالت پیوسته علامت جمع در مخرج رابطه (۷.۵) با علامت انتگرال تعویض می‌شود. به دلیل وجود $\sum \exp\{Q(Z)\}$ و $\int \exp\{Q(Z)\}$ در روابط بالا در اکثر توزیع‌ها توزیع توام Z صورت بسته ندارد و برای ماکسیمم کردن تابع درست‌نمایی نیاز به روش‌های عددی است. دقت کنید که میزان اطلاع درباره $Q(\cdot)$ هم‌ارز میزان اطلاع درباره $P(\cdot)$ است.

قضیه ۲.۴.۵. (بیسگ، ۱۹۷۴). تابع $Q(\cdot)$ دارای دو ویژگی زیر است

۱. تابع $Q(\cdot)$ در رابطه

$$\frac{P(Z(s_i)|Z_{-i})}{P(\circ(s_i)|Z_{-i})} = \frac{P(Z)}{P(\circ)} = e^{Q(Z)-Q(Z_i)} \quad (۸.۵)$$

صدق می‌کند، که در آن $\circ(s_i)$ نشانگر پیشامد $\circ$ است، $Z(s_i)$ است،

$$Z_{-i} \equiv \{Z(s_j); j \neq i\}$$

$$\text{و } Z_i \equiv (Z(s_1), \dots, Z(s_{i-1}), \circ, Z(s_{i+1}), \dots, Z(s_n))$$

۲. تابع $Q(\cdot)$ به صورت یکتا روی تکیه‌گاه Z ها به صورت

$$\begin{aligned} Q(Z) &= \sum_{1 \leq i \leq n} Z(s_i) G_i(Z(s_i)) \\ &+ \sum_{1 \leq i \leq j \leq n} Z(s_i) Z(s_j) G_{ij}(Z(s_i), Z(s_j)) \\ &\vdots \\ &+ Z(s_1) \cdots Z(s_n) G_{1, \dots, n}(Z(s_1), \dots, Z(s_n)) \end{aligned} \quad (9.5)$$

بسط داده می‌شود، که در آن $G(\cdot)$ تابع پیوند است. در واقع $Q(\cdot)$ می‌تواند برحسب مجموع توابعی از زیرمجموعه‌های $Z(s_1), \dots, Z(s_n)$ نمایش داده شود، که این زیرمجموعه‌ها توسط جرگه‌ها مشخص می‌شوند.

قضیه ۳.۴.۵. (هامرسلی - کلیفورد، ۱۹۷۱). فرض کنید Z یک میدان تصادفی مارکوفی باشد، که در شرط مثبت بودن صدق می‌کند، در این صورت شرط لازم برای بسط (۹.۵) آن است که اگر مکان‌های $s, \dots, j, i$ در یک جرگه نباشند، آن‌گاه، $G_{ij \dots s}(\cdot) = 0$ باشد. قضیه هامرسلی - کلیفورد باعث شده است که میدان‌های تصادفی مارکوفی کاربردهای بسیار زیادی در متون فضایی پیدا کنند.

۵.۵ توزیع‌های نمایی شرطی

در این بخش خانواده توزیع‌های نمایی برای مدل‌بندی توزیع‌های شرطی

$$P(Z(s_i) | Z(s_j), j \in N_i)$$

را معرفی می‌کنیم. این خانواده منعطف، بسیاری از توزیع‌ها مانند توزیع دوجمله‌ای، پواسون، گاوسی و گاما را در بر می‌گیرد. همچنین مدل‌های فضایی که شکل نمایی دارند، به راحتی تفسیر می‌شوند (بیسگ، ۱۹۷۴). در حالت گسسته خانواده توزیع‌های نمایی شرطی بر اساس پارامترهای طبیعی به صورت

$$P(Z(s_i) | Z_{-i}) = e^{L_i(Z_{-i}) M_i(Z(s_i)) + N_i(Z(s_i)) + V_i(Z_{-i})} \quad (10.5)$$

نوشته می‌شوند که در آن $L_i(\cdot)$ و $M_i(\cdot)$ ضابطه‌های معلومی دارند و $N_i(\cdot)$ ، $V_i(\cdot)$ توابعی از مقادیر مشاهده‌شده در همسایگی‌های موقعیت i هستند. بیسگ (۱۹۷۴) برای خانواده توزیع‌های نمایی شرطی با فرض وابستگی فقط جفتی بین موقعیت‌ها (یعنی $G_a(\cdot) = 0$ برای هر a که تعداد عناصر آن سه یا بیشتر باشد) رابطه زیر را اثبات کرد. در حقیقت این نتیجه از قضیه هامرسلی - کلیفورد نتیجه می‌شود:

$$A_i(Z_{-i}) = \alpha_i + \sum_{j=1}^n \theta_{ij} M_j(Z(s_j)) \quad j = 1, 2, \dots, n, j \neq i \quad (11.5)$$

که در آن $\theta_{ij} = G_{ij}(Z(s_i), Z(s_j))$ و $\theta_{ij} = \theta_{ji}$. بنابراین در صورتی که $j \notin N_i$ (همسایه نبودن i و j)، $\theta_{ij} = 0$. همچنین برای شناسایی پذیر بودن مدل شرط $\theta_{ii} = 0$ در نظر گرفته می‌شود. با فرض وابستگی فقط جفتی، رده‌ای از مدل‌ها برای داده‌های گسسته مانند اتولجستیک، اتودوجمله‌ای و اتوپواسون و برای داده‌های پیوسته مدل اتوگاوسی معرفی شده‌اند. کاربردهایی از این مدل‌ها را می‌توان در مدل‌بندی میزان بروز ناحیه‌ای بیماری‌های خاص در پزشکی (کاوسی و همکاران، ۱۳۸۵؛ لاوسون، ۲۰۰۱)، مکان‌یابی مناسب برای ساخت و ساز مسکن در شهرسازی، بازاریابی و تعیین مکان‌های مناسب برای ایجاد فروشگاه‌های زنجیره‌ای نام برد. در ادامه مدل‌های معمول برای پاسخ‌های گسسته را معرفی می‌کنیم.

اتولجستیک

فرض کنید متغیر مورد مطالعه دودویی صفر و یک باشد، که وجود یا عدم وجود یک مشخصه را بیان می‌کند. تحت پذیره وابستگی فقط جفتی بین موقعیت‌ها، توزیع شرطی متغیر دودویی $Z(s_i)$ لجستیک است. با قرار دادن $G_i(1) \equiv \alpha_i$ و $G_{ij}(1, 1) = \theta_{ij}$ در (۹.۵) می‌توان نوشت

$$Q(Z) = \sum_{i=1}^n \alpha_i Z(s_i) + \sum_{1 \leq i < j \leq n} \theta_{ij} Z(s_i) Z(s_j)$$

با محاسبه

$$Q(Z) - Q(Z_i) = \alpha_i Z(s_i) + \sum_{j=1}^n \theta_{ij} Z(s_i) Z(s_j)$$

و جای‌گذاری در (۸.۵) و استفاده از برابری

$$P(Z(s_i) = 1 | Z_{-i}) = 1 - P(0(s_i) | Z_{-i})$$

و در نظر گرفتن این که $Z(s_i) = 1$ یا $Z(s_i) = 0$ ، نتیجه می‌شود

$$\frac{1 - P(0(s_i) | Z_{-i})}{P(0(s_i) | Z_{-i})} = \exp\{\alpha_i + \sum_{j=1}^n \theta_{ij} Z(s_j)\}.$$

بنابراین

$$\frac{1}{P(0(s_i) | \{z(s_j) : j \neq i\})} = 1 + \exp\{\alpha_i + \sum_{j=1}^n \theta_{ij} z(s_j)\}.$$

پس

$$P(0(s_i) | \{z(s_j) : j \neq i\}) = \frac{1}{1 + \exp\{\alpha_i + \sum_{j=1}^n \theta_{ij} z(s_j)\}}$$

یا معادل آن

$$\begin{aligned} \text{logit}(P(0(s_i) | Z_{-i})) &= \ln \left(\frac{1 - P(0(s_i) | Z_{-i})}{P(0(s_i) | Z_{-i})} \right) \\ &= \alpha_i + \sum_{j=1}^n \theta_{ij} Z(s_j) \end{aligned}$$

نتیجه می‌شود. در نتیجه مدل اتولجستیک به صورت

$$P(Z(s_i)|Z_{-i}) = \frac{\exp\{\alpha_i Z(s_i) + \sum_{j=1}^n \theta_{ij} Z(s_i) Z(s_j)\}}{1 + \exp\{\alpha_i Z(s_i) + \sum_{j=1}^n \theta_{ij} Z(s_i) Z(s_j)\}}$$

به دست می‌آید، که در آن پارامترهای روند فضایی (تغییرات بزرگ مقیاس) و θ_{ij} پارامترهای وابستگی فضایی (تغییرات کوچک مقیاس) هستند.

مدل اتودوجمله‌ای

مدل اتودوجمله‌ای به صورت

$$\begin{aligned} P(Z(s_i) | \{z(s_j) : j \neq i\}) \\ \equiv \binom{n_i}{z(s_i)} \pi_i(\{z(s_j) : j \neq i\})^{z(s_i)} (1 - \pi_i(\{z(s_j) : j \neq i\}))^{n_i - z(s_i)} \quad z(s_i) = 0, 1, \dots, n \end{aligned} \quad (12.5)$$

تعریف می‌شود، که صورت نمایش نمایی آن به شکل

$$P(Z(s_i) | Z_{-i}) = \exp^{Z(s_i) \logit(\pi_i) + \ln \binom{n_i}{z(s_i)} + n_i \ln(1 - \pi_i)}$$

است که در آن $L_i(Z_{-i}) = \logit(\pi_i) = \ln \frac{\pi_i}{1 - \pi_i}$ ، $\pi = \pi(Z_{-i})$ ، $M_i(z(s_i)) = z(s_i)$ و $M_i(z(s_i)) = z(s_i)$ با توجه به تعریف داریم

$$\begin{aligned} Q(Z) &= \ln \left(\frac{P(z(s_i) | \{z(s_j) : j \neq i\})}{P(0(s_i) | \{z(s_j) : j \neq i\})} \right) \\ &= \ln \left(\frac{\binom{n_i}{z(s_i)} \pi_i^{z(s_i)} (1 - \pi_i)^{n_i - z(s_i)}}{(1 - \pi_i)^{n_i}} \right) \\ &= \ln \left(\binom{n_i}{z(s_i)} \pi_i^{z(s_i)} (1 - \pi_i)^{-z(s_i)} \right) \\ &= \ln \left(\binom{n_i}{z(s_i)} \left(\frac{\pi_i}{1 - \pi_i} \right)^{z(s_i)} \right) \\ &= z(s_i) \left\{ \ln \left(\frac{\pi_i}{1 - \pi_i} \right) + \ln \binom{n_i}{z(s_i)} \right\} \\ &= \alpha_i z(s_i) + \sum_{i=1}^n \theta_{ij} z(s_i) z(s_j) + \ln \binom{n_i}{z(s_i)} \\ Q(Z) &= \sum_{i=1}^n \alpha_i z(s_i) + \sum_{n \leq j \leq i \leq n} \sum \theta_{ij} z(s_i) z(s_j) + \sum_{i=1}^n \ln \left(\binom{n_i}{z(s_i)} \right). \end{aligned}$$

بنابراین معادله

$$\logit(\pi_i) = \alpha_i + \sum_{j=1}^n \theta_{ij} z(s_j) \quad i = 1, \dots, n$$

که همان مدل اتولجستیک است، به دست می‌آید. با حل آن بر حسب π_i داریم

$$\pi_i(\{z(s_j) : j \neq i\}) = \frac{\exp\{\alpha_i + \sum_{j=1}^n \theta_{ij} z(s_j)\}}{1 + \exp\{\alpha_i + \sum_{j=1}^n \theta_{ij} z(s_j)\}} \quad (۱۳.۵)$$

که در آن θ_{ij} ها همبستگی فضایی و α_i تاثیر موقعیت i ام است به طوری که می‌تواند به عنوان روند مدل یا تاثیر متغیرهای کمکی در نظر گرفته شود. از رابطه (۱۲.۵) می‌توان به عنوان پیش‌گو و رابطه (۱۳.۵) به عنوان برآوردگر در موقعیت i ام استفاده کرد.

مدل اتوپواسون

توزیع شرطی اتوپواسون با پذیره وابستگی فقط جفتی به صورت

$$P(Z(s_i) | \mathbf{Z}_{-i}) = \frac{e^{-\lambda_i(\mathbf{Z}_{-i})} \{\lambda_i(\mathbf{Z}_{-i})\}}{Z(s_i)!}$$

است. برای این مدل می‌توان نوشت

$$\begin{aligned} Q(\mathbf{Z}) - Q(\mathbf{Z}_i) &= \ln \left(\frac{P(z(s_i) | z(s_j) : j \neq i)}{P(\circ(s_i) | z(s_j) : j \neq i)} \right) \\ &= \ln \frac{\frac{e^{-\lambda_i Z(s_i)} \lambda_i^{Z(s_i)}}{Z(s_i)!}}{e^{-\lambda_i Z(s_i)}} = \ln \frac{\lambda_i^{Z(s_i)}}{Z(s_i)!} \\ &= Z(s_i) \ln(\lambda_i) - \ln(Z(s_i)!). \end{aligned}$$

با توجه به (۱۱.۵) خواهیم داشت

$$\begin{aligned} Q(\mathbf{Z}) &= Z(s_i) \{\alpha_i + \sum_{j=1}^n \theta_{ij} Z(s_j)\} - \ln(Z(s_i)!) \\ &= \alpha_i Z(s_i) - \ln(Z(s_i)!) + \sum_{j \neq i} \theta_{ij} Z(s_i) Z(s_j). \end{aligned}$$

بنابراین

$$Q(\mathbf{Z}) = \sum_{i=1}^n \alpha_i Z(s_i) + \sum_{1 \leq i \leq j \leq n} \theta_{ij} Z(s_i) Z(s_j) - \sum_{i=1}^n \ln(Z(s_i)!).$$

مدل اتودوجمله‌ای منفی

مدل فضایی دوجمله‌ای منفی با پارامترهای میانگین μ و بیش پراکندگی β ، تحت عنوان اتودوجمله‌ای منفی، برای داده‌های شمارشی به صورت

$$P(Z(s_i) | \mathbf{Z}_{-i}) = \frac{\Gamma(Z(s_i) + \mu_i)}{Z(s_i)! \Gamma(\mu_i)} \left(\frac{\beta_i \{\mathbf{Z}_{-i}\}}{\beta_i \{\mathbf{Z}_{-i}\} + 1} \right)^{\mu_i} \left(\frac{1}{\beta_i \{\mathbf{Z}_{-i}\} + 1} \right)^{Z(s_i)} \quad Z(s_i) = 0, 1, \dots$$

تعریف می‌شود. در صورت وابستگی فقط جفتی بین متغیرها

$$\beta_i \{\mathbf{Z}_{-i}\} = e^{-\alpha_i - \sum_{j=1}^n \theta_{ij} Z(s_j)} - 1$$

است.

در ادامه برای معرفی مدل اتوبتا-دوجمله‌ای، ابتدا نسخه بازپارامتری شده آن را معرفی می‌کنیم.

۶.۵ توزیع بتا-دوجمله‌ای بازپارامتری شده

فرض کنید

$$Z(s_i) | \pi(s_i) \sim \text{bin}(n(s_i), \pi(s_i))$$

به‌طوری که

$$\pi(s_i) \sim \text{Beta}(a(s_i), b(s_i)).$$

بنابراین $Z(s_i)$ دارای توزیع کناری بتا-دوجمله‌ای با تکیه‌گاه $\{0, \dots, n_s\}$ و تابع احتمال

$$P(Z_i(s) = z(s_i) | a(s_i), b(s_i)) = \binom{n_s}{z(s_i)} \frac{B(a(s_i) + z(s_i), n_s + b(s_i) - z(s_i))}{B(a(s_i), b(s_i))}$$

می‌باشد. میانگین و واریانس $Z_i(s)$ به‌ترتیب برابر هستند با $n_s \frac{a(s_i)}{a(s_i) + b(s_i)}$ و

$$\frac{n_s a(s_i) b(s_i) (n_s + a(s_i) + b(s_i))}{(a(s_i) + b(s_i))^2 (1 + a(s_i) + b(s_i))}.$$

حال میانگین و واریانس توزیع بتا-دوجمله‌ای را بازپارامتری می‌کنیم.

فراری و سریباری (۲۰۰۴) توزیع بتای بازپارامتری را برای مطالعه نسبت‌ها پیشنهاد کردند. آن‌ها پارامترهای توزیع بتا را به‌گونه‌ای بازپارامتری کردند که مدل رگرسیونی براساس میانگین متغیر پاسخ تعیین شود.

فرض کنید متغیر پاسخ Z_i دارای توزیع بتا با پارامترهای $a(s_i)$ و $b(s_i)$ باشد، که در آن $a(s_i), b(s_i) > 0$. در این‌صورت، میانگین توزیع برابر $\mu(s_i) = \frac{a(s_i)}{a(s_i) + b(s_i)}$ خواهد بود. با قرار دادن $a(s_i) = \gamma(s_i) \mu(s_i)$ و $b(s_i) = \gamma(s_i) (1 - \mu(s_i))$ پارامترهای اصلی به‌صورت

$$b(s_i) = \gamma(s_i) (1 - \mu(s_i))$$

می‌توانند بازنویسی شوند. حال با توجه به تعاریف فوق تابع چگالی توزیع بتا را به‌صورت زیر بازنویسی می‌کنیم:

$$f_i(z_i | \mu(s_i), \gamma(s_i)) = \frac{\Gamma(\gamma(s_i))}{\Gamma(\mu(s_i)\gamma(s_i))\Gamma((1 - \mu(s_i))\gamma(s_i))} z_i^{\mu(s_i)\gamma(s_i)-1} (1 - z_i)^{(1 - \mu(s_i))\gamma(s_i)-1} I_z(0, 1)$$

$$0 < \mu(s_i) < 1, \gamma(s_i) > 0$$

در این‌صورت

$$Z_i \sim \text{Beta}(\mu(s_i)\gamma(s_i), (1 - \mu(s_i))\gamma(s_i))$$

با میانگین μ_i و واریانس $\sigma^2 = \frac{\mu(s_i)(1-\mu(s_i))}{\gamma(s_i)+1}$ است. با استفاده از ویژگی‌های رگرسیون بتا و بازپارامتری کردن توزیع بتا-دوجمله‌ای با عبارات

$$\mu(s_i) = \frac{a(s_i)}{\gamma(s_i)}$$

$$\gamma(s_i) = a(s_i) + b(s_i)$$

تابع احتمال $Z(s_i)$ را می‌توان به صورت

$$Z(s_i) \sim \text{beta-bin}(n_s, \mu(s_i)\gamma(s_i), (1 - \mu(s_i))\gamma(s_i))$$

نوشت. میانگین و واریانس توزیع بتا-دوجمله‌ای بازپارامتری نیز به صورت

$$E(Z(s_i)) = n_s \mu(s_i)$$

و

$$\text{Var}(Z(s_i)) = n_s \mu(s_i)(1 - \mu(s_i)) \frac{\gamma(s_i) + n_s}{\gamma(s_i) + 1}$$

به دست می‌آیند. لازم به ذکر است پارامتر بیش‌پراکندگی برای مدل بتا-دوجمله‌ای توسط

$$\frac{\gamma(s) + n_s}{\gamma(s) + 1} \in (1, n_s)$$

لحاظ می‌شود. همچنین زمانی که $n_s = 1$ یا $\gamma(s) \rightarrow \infty$ واریانس کناری توزیع بتا-دوجمله‌ای به واریانس دوجمله‌ای نزدیک می‌شود.

۱.۶.۵ مدل اتوبتا-دوجمله‌ای

در این بخش به تعمیم مدل بتا-دوجمله‌ای به اتوبتا-دوجمله‌ای می‌پردازیم. مدل اتوبتا-دوجمله‌ای به صورت

$$P(Z(s_i) | \{z(s_j) : j \neq i\}) = \binom{n_s}{z(s_i)} \frac{B(a(s_i) + z(s_i), n_s + b(s_i) - z(s_i))}{B(a(s_i), b(s_i))}$$

تعریف می‌شود. با توجه به (۸.۵)

$$Q(\mathbf{Z}) = \log \left(\frac{P(Z(s_i) | \{z(s_j) : j \neq i\})}{P(z(s_i) | \{z(s_j) : j \neq i\})} \right) = \log \left(\binom{n_s}{z(s_i)} \frac{B(a(s_i) + z(s_i), n_s + b(s_i) - z(s_i))}{B(a(s_i), b(s_i) + n_s)} \right)$$

با لگاریتم گرفتن، داریم

$$\begin{aligned} Q(\mathbf{Z}) &= \log B(a(s_i) + z(s_i), n_s + b(s_i) - z(s_i)) + \log \binom{n_s}{z(s_i)} - \log B(a(s_i), b(s_i)) \\ &= \text{lgamma}(a(s_i) + z(s_i)) + \text{lgamma}(n_s + b(s_i) - z(s_i)) + \log \binom{n_s}{z(s_i) + n_s} \\ &\quad - \text{lgamma}(a(s_i) + b(s_i) + n) - \log B(a(s_i), b(s_i)) \end{aligned}$$

که در آن منظور از $lgamma$ ، لگاریتم تابع گاما است. در نهایت تابع احتمال مدل اتوبتا-دوجمله‌ای به صورت زیر حاصل می‌شود:

$$P(Z(s_i) | \{z(s_j) : j \neq i\}) = \binom{n_i}{z(s_i)} \frac{B(\mu(s_i)\gamma(s_i) + z_i, n_s + 1 - \mu(s_i)\gamma(s_i) - z_i)}{B(\mu(s_i)\gamma(s_i), (1 - \mu(s_i)\gamma(s_i)))}$$

۷.۵ استنباط بیزی تقریبی

گاهی در استنباط بیزی توزیع پسین مدل‌ها شکل بسته‌ای ندارند. در چنین مواردی معمولاً از الگوریتم‌های نمونه‌گیری مونت کارلویی برای تقریب توزیع استفاده می‌شود. وجود همبستگی در میدان تصادفی پنهان معمولاً موجب کاهش کارایی این الگوریتم‌ها، افزایش زمان محاسبات و همگرایی کند الگوریتم می‌شود. برای مرتفع کردن این مشکل، رو و همکاران (۲۰۰۹) روش تقریب لاپلاس آشیانه‌ای جمع‌بسته را معرفی کردند، که در استنباط مدل‌های گاوسی پنهان جانشین مناسبی برای الگوریتم‌های MCMC است. در روش INLA انتگرال‌گیری عددی و تقریب لاپلاس به‌طرقی کارا ترکیب می‌شوند به‌طوری که محاسباتی سریع و با دقت قابل قبول جایگزین شبیه‌سازی‌های سنگین MCMC می‌شوند.

۱.۷.۵ مدل‌های گاوسی پنهان

گام اول در تعریف یک مدل گاوسی پنهان، همراه با چارچوب بیزی، تشخیص توزیع داده‌های مشاهده‌شده $Z = (Z_1, \dots, Z_n)$ است. پرکاربردترین روش تشخیص توزیع Z_i با پارامتر μ_i (معمولاً میانگین $E(Z_i)$)، تابع پیش‌گوی جمعی η_i از طریق تابع پیوند $g(\cdot)$ ، مانند $g(\mu_i) = \eta_i$ است. پیش‌گوی خطی جمعی η_i به‌صورت زیر تعریف می‌شود:

$$\eta_i = \alpha + \sum_{m=1}^M \beta_m v_{mi} + \sum_{l=1}^L f_l(t_{li}) + \epsilon_i \quad i = 1, \dots, n \quad (14.5)$$

که در آن α عرض از مبدا است، ضرایب $\beta = \{\beta_1, \dots, \beta_M\}$ اثرات ثابت خطی متغیرهای تبیینی $v = (v_1, \dots, v_M)$ روی متغیر پاسخ است و $f = \{f_1(\cdot), \dots, f_L(\cdot)\}$ مجموعه‌ای از توابع است که در جملات بردار متغیرهای تصادفی و پنهان $t = (t_1, \dots, t_L)$ تعریف شده‌اند. همه مولفه‌های پنهان (غیرقابل مشاهده) مورد نظر برای استنباط در مجموعه پارامترها با نام x که به‌صورت $x = \{\alpha, \beta, f\}$ تعریف شده، جمع‌آوری می‌شود. به‌علاوه بردار m بعدی ابرپارامترها به‌صورت $\theta = \{\theta_1, \dots, \theta_m\}$ مشخص می‌شود. با فرض استقلال شرطی، توزیع شرطی n مشاهده به‌وسیله تابع درستنمایی زیر به‌دست می‌آید:

$$\pi(z|x, \theta) = \prod_{i=1}^n \pi(z_i|x_i, \theta). \quad (15.5)$$

فرض می‌کنیم پیشین x دارای توزیع نرمال چندمتغیره با بردار میانگین $\circ$ و ماتریس دقت

$$Q(\theta) = \Sigma_{\theta}^{-1}$$

یعنی $x \sim N(\circ, Q(\theta))$ با تابع چگالی

$$\pi(x|\theta) = (\pi)^{-\frac{n}{2}} |Q(\theta)|^{\frac{1}{2}} \exp\left(-\frac{1}{2} x' Q(\theta) x\right) \quad (۱۶.۵)$$

باشد. در حالتی که $Q(\theta)$ یک ماتریس دقت تنک^{۱۴} باشد، مولفه‌های میدان گاوسی پنهان x ، به‌طور شرطی، مستقل هستند. این ویژگی با عنوان میدان تصادفی مارکوف گاوسی^{۱۵} (GMRF) شناخته می‌شود (رو و هلد، ۲۰۰۵). توزیع توام x و θ به‌وسیله ضرب درستنمایی (۱۵.۵) و چگالی (۱۶.۵) و توزیع پیشین ابرپارامترهای $\pi(\theta)$ به دست می‌آید:

$$\pi(x, \theta|z) \propto \pi(\theta) \pi(x|\theta) \pi(z|x, \theta) \quad (۱۷.۵)$$

$$\begin{aligned} &\propto \pi(\theta) \pi(x|\theta) \prod_{i=1}^n \pi(z_i|x_i, \theta) \\ &\propto \pi(\theta) \times |Q(\theta)|^{-\frac{n}{2}} \exp\left(-\frac{1}{2} x' Q(\theta) x\right) \times \prod_{i=1}^n \exp(\log(\pi(z_i|x_i, \theta))) \\ &\propto \pi(\theta) |Q(\theta)|^{\frac{n}{2}} \exp\left(-\frac{1}{2} x' Q(\theta) x + \sum_{i=1}^n \log(\pi(z_i|x_i, \theta))\right). \end{aligned}$$

۲.۷.۵ تقریب لاپلاس آشیانه‌ای جمع‌بسته

برای تحلیل بیزی مدل‌های گاوسی پنهان به روش INLA، لازم است توزیع‌های پسینی کناری متغیرهای پنهان و ابرپارامترها به‌صورت

$$\pi(x_i|z) = \int \pi(x_i|\theta, z) \pi(\theta|z) d\theta \quad i = 1, \dots, n_d \quad (۱۸.۵)$$

$$\pi(\theta_j|z) = \int \pi(\theta|z) d\theta_{-j} \quad j = 1, \dots, m \quad (۱۹.۵)$$

محاسبه شوند، که در آن بردار حاصل از حذف درایه j ام θ است و n_d بعد میدان تصادفی x است. روش INLA تقریب‌هایی برای چگالی‌های پسین (۱۸.۵) و (۱۹.۵) به‌صورت

$$\tilde{\pi}(x_i|z) = \int \tilde{\pi}(x_i, z, \theta) \tilde{\pi}(\theta|z) d\theta, \quad i = 1, \dots, n$$

$$\tilde{\pi}(\theta_j|z) = \int \tilde{\pi}(\theta|z) d\theta_{-j} \quad j = 1, \dots, m$$

^{۱۴} Sparse

^{۱۵} Gaussian Markov random field

ارایه می‌کند. روش INLA برای محاسبه تقریبی این توزیع‌ها، با استفاده از تبدیل‌های لاپلاس و انتگرال‌گیری عددی، محاسباتی سریع و تقریبی دقیق را جایگزین شبیه‌سازی‌های سنگین الگوریتم‌های MCMC می‌کند. تقریب‌های $\tilde{\pi}(x_i|z)$ و $\tilde{\pi}(\theta_i|z)$ در سه گام به شرح زیر محاسبه می‌شوند:

۱. توزیع پسینی کناری $\pi(\theta|z)$ به صورت

$$\tilde{\pi}(\theta|z) = \frac{\pi(x, \theta, z)}{\tilde{\pi}(x|\theta, z)} \Big|_{x=x^*(\theta)}$$

تقریب زده می‌شود، که در آن $\tilde{\pi}_G(x|\theta, z)$ تقریب گاوسی (رو و هلد، ۲۰۰۵) برای توزیع شرطی کامل x و x^* مد توزیع شرطی کامل x است.

۲. محاسبه $\pi(x_i|\theta, z)$ که توسط رو همکاران (۲۰۰۹) سه تقریب گاوسی، لاپلاس و لاپلاس ساده‌شده^{۱۶} را برای تقریب پیشنهاد دادند. ساده‌ترین روش، تقریب گاوسی و دقیق‌ترین روش، تقریب لاپلاس است. در نهایت، روش تقریب لاپلاس ساده شده با کاهش جزئی در دقت از لحاظ محاسباتی آسان‌تر از تقریب لاپلاس است.

۳. گام سوم ترکیب دو گام پیشین و استفاده از روش انتگرال‌گیری عددی برای رسیدن به هدف است. توزیع پسین به صورت

$$\tilde{\pi}(x_i|z) = \sum_{k=1}^K \tilde{\pi}(x_i|\theta_k, z) \times \tilde{\pi}(\theta_k|z) \times \Delta_k$$

به دست می‌آید که در آن K تعداد θ های انتخاب شده از تکیه‌گاه و Δ_k ها وزن‌های این نقاط هستند.

در ادامه از تقریب INLA برای تحلیل داده‌های واقعی سرطان معده بر اساس مدل اتوبتا-دوجمله‌ای استفاده می‌کنیم.

۸.۵ تحلیل داده‌های سرطان معده در استان گیلان

همزمان با رشد روزافزون اطلاعات پیرامون بیماری‌ها و مرگ و میر، روش‌های متناسب برای تحلیل این نوع داده‌ها که پاسخ‌گوی نیازهای مختلف باشد نیز رو به گسترش است. یکی از این روش‌ها، پهنه‌بندی بیماری یا مرگ و میر است که توزیع جغرافیایی بیماری‌ها یا مرگ را در کنار دیگر عوامل خطر در نظر می‌گیرد. پهنه‌بندی بیماری یا مرگ و میر به مجموعه‌ای از روش‌های آماری اطلاق می‌شود که هدف آن‌ها به دست آوردن برآوردهایی دقیق از میزان بروز یا شیوع بیماری‌ها یا مرگ و میر و تنظیم آن‌ها در قالب نقشه‌های جغرافیایی می‌باشد.

^{۱۶} Simplified laplace

امروزه پهنه‌بندی و برآورد خطر بیماری‌ها مورد توجه فعالان و برنامه‌ریزان عرصه سلامت جامعه می‌باشد چرا که توزیع جغرافیایی میزان‌های بروز، شیوع و مرگ و میر نقش مهمی در تخصیص عوامل خطر و پیش‌گیری از آن‌ها بازی می‌کند. تحلیل جغرافیایی نرخ‌های بیماری علاوه بر فرمول‌بندی و ارزیابی پذیره‌های سبب‌شناختی و اعمال مداخله در مناطقی که نیازمند توجه خاص هستند، می‌تواند نقش مهمی در زمینه تخصیص منابع، امکانات و نیروی انسانی ایفا کند.

شرایط جوی و محیطی در هر منطقه زمینه را برای بروز و شیوع برخی بیماری‌ها مساعد می‌کند. سرطان نیز از جمله بیماری‌های کشنده و نگران‌کننده امروزی است که انسان با آن دست و پنجه نرم می‌کند. عوامل عمده تاثیرگذار بر این بیماری را عوامل محیطی و ژنتیکی می‌دانند. میزان بروز انواع آن در نواحی جغرافیایی مختلف، متفاوت است. در سال‌های اخیر در کشور ما این میزان افزایش پیدا کرده است. به‌خصوص در استان گیلان، سرطان معده شیوع فراوان دارد و سالانه جان صدها نفر را به خطر می‌اندازد. استان گیلان از نظر فراوانی سرطان معده مقام اول را در کل کشور دارد (رمضانی، ۱۳۸۷).

نسبت بالایی از سرطان‌ها به‌طور مستقیم با عوامل محیط زیست ارتباط دارند و لازم است که عوامل موثر در این مورد تعیین و مشخص گردد (ویلکاکسون و تنسر، ۱۹۹۹). از هر ۹ مرگ در جهان یک مورد مربوط به سرطان است. سرطان بعد از بیماری‌های قلبی عروقی دومین علت مرگ در جوامع است و سومین علت مرگ بعد از بیماری‌های قلبی عروقی و تصادفات در کشور ما است. با توجه به رابطه بین بروز انواع سرطان و شرایط جغرافیایی منطقه‌ای، به همین علت تفاوت‌های آشکاری را در میزان شیوع و فراوانی هر کدام از سرطان‌ها در مناطق مختلف شاهد هستیم به‌طوری که در عرض بالای ۳۲ درجه جغرافیایی شیوع بیماری سرطان مری بالا است و تاثیر زمینه‌های محیطی و شرایط جغرافیایی و نوع اشتغال را با این بیماری در ارتباط دانسته‌اند.

سرطان معده عبارت است از رشد بدون کنترل سلول‌های بدخیم در معده که در آن بیش‌تر افراد تا مراحل پیشرفته بیماری، علامتی ندارند. سلول‌های سرطانی معمولاً به مرور رشد می‌کنند و تبدیل به تومور می‌شوند. این سرطان در استان گیلان بسیار شایع است، به‌طوری که به گفته یکی از مسئولان مرکز تحقیقات گوارش و کبد بیمارستان رازی رشت، هر دو روز یک مورد سرطان جدید معده تشخیص داده می‌شود و از هر سه نفری که بر اثر ابتلا به سرطان فوت می‌شوند یک نفر مبتلا به سرطان معده است.

در مورد روند بهبود برآوردهای پهنه‌بندی بیماری می‌توان به تحقیقات انجام‌شده اشاراتی داشت. ابتدایی‌ترین نقشه را می‌توان به جان اسنو نسبت داد که در سال ۱۸۵۴ و به دنبال همه‌گیری وبا در شهر لندن تهیه کرد. کلایتون و کالدور (۱۹۸۷) مدل‌های سلسله‌مراتبی و استنباط بیز تجربی مرتبط با آن را برای میزان‌های مرگ و میر استانداردشده در حالی که همبستگی فضایی بین مشاهدات در نواحی همسایه نیز لحاظ می‌شود، مطرح کردند.

در تحلیل داده‌های بیماری که غالباً گسسته و شمارشی هستند، می‌توان کایسر و همکارانش

(۱۹۹۷) را نام برد که برای تحلیل فضایی داده‌های دارای توزیع پواسون، مدل توزیع پواسون فضایی را معرفی نمودند. الیور و همکارانش (۱۹۹۸) هم‌کریگینگ دوجمله‌ای را برای تهیه پهنه‌بندی خطر سرطان کودکان در غرب انگلستان به کار بردند. مونستیز و همکارانش (۲۰۰۴-۲۰۰۶) استفاده از مدل کریگینگ پواسون را برای داده‌های شمارشی همبسته فضایی پیشنهاد دادند. آن‌ها در مدل پیشنهادی خود از یک طرح وزن دار به منظور اختصاص دادن وزن بزرگ‌تر به مشاهدات بزرگ‌تر استفاده کردند.

۱.۸.۵ نرخ سرطان معده در استان گیلان

در این بخش نرخ سرطان معده در استان گیلان را با روش INLA، تحلیل می‌کنیم. در این راستا دو هدف را دنبال می‌کنیم: (۱) تهیه یک نقشه پهنه‌بندی از نرخ سرطان معده در استان گیلان و (۲) بررسی تاثیر متغیرهای تبیینی.

مجموعه داده مورد استفاده شامل کلیه داده‌های ثبت‌شده بیماران سرطانی در مرکز آموزشی درمانی رازی رشت از اول خرداد ۱۳۹۱ تا پایان اسفند ۱۳۹۵ می‌باشد. کل بیماران مبتلا به سرطان معده در این مجموعه داده‌ها ۶۴۴ مورد هستند که ۶۲ درصد آن‌ها مرد و ۳۸ درصد زن هستند. در سال ۹۱ کم‌ترین مراجعه صورت گرفته و بیش‌ترین مراجعه نیز در سال ۹۵ است که دلیل آن شاید به‌خاطر ثبت دقیق‌تر اطلاعات برای سال ۹۵ نسبت به سال‌های دیگر باشد. متغیر پاسخ تعداد افراد بیمار در شهرهای مختلف استان گیلان و متغیرهای تبیینی شامل سرانه فضای سبز و مدت اقامت بیماران (تعداد روزهای بستری) هستند. با توجه به این‌که متغیر پاسخ شمارشی و گسسته است و مشاهدات به موقعیت جغرافیایی که در آن قرار دارند وابسته هستند، بنابراین داده‌ها ماهیت فضایی و شبکه‌ای دارند. برای تحلیل داده‌های سرطانی با استفاده از مدل بتا-دوجمله‌ای و دوجمله‌ای، پیش‌گوی خطی

$$\eta_i = \beta_0 + \beta_1 x_{1i} + \beta_2 x_{2i} + u_i, \quad i = 1, \dots, n$$

را در نظر گرفتیم، که در آن $\beta = (\beta_0, \beta_1, \beta_2)'$ بردار ضرایب رگرسیونی، x_1 نرخ بستری بیمار در بیمارستان و x_2 نرخ سرانه فضای سبز است. همچنین $u = (u_1, \dots, u_{16})'$ اثر تصادفی فضایی CAR برای لحاظ کردن همبستگی فضایی بین شانزده شهرستان استان گیلان، وارد مدل شده است. در مدل CAR در نظر گرفته‌شده برای تحلیل این داده‌ها از یک مدل ذاتی^{۱۷} استفاده کردیم که در آن فرض می‌شود $\rho = 1$ است و تنها ابرپارامتر این مدل پارامتر واریانس σ^2 است. البته در روش INLA به جای کار کردن با σ^2 و تعریف توزیع پیشین برای آن، با عکس آن که پارامتر دقت محسوب می‌شود مدل‌بندی (و تعریف توزیع پیشین) انجام می‌شود.

برای تکمیل مدل‌بندی بیزی، باید توزیع‌های پیشین پارامترهای β و θ تعیین شوند. برای این کار، در هر دو مدل دوجمله‌ای و بتا-دوجمله‌ای، از توزیع‌های پیشین استاندارد روش

^{۱۷}Intrinsic

INLA استفاده کردیم؛ برای پارامترهای β توزیع‌های پیشین نرمال با میانگین صفر و واریانس بزرگ ۱۰۰۰ تعیین شدند و برای پارامتر دقت مدل CAR توزیع پیشین گاما با پارامترهای شکل ۱ و مقیاس ۰/۰۰۱ انتخاب شد.

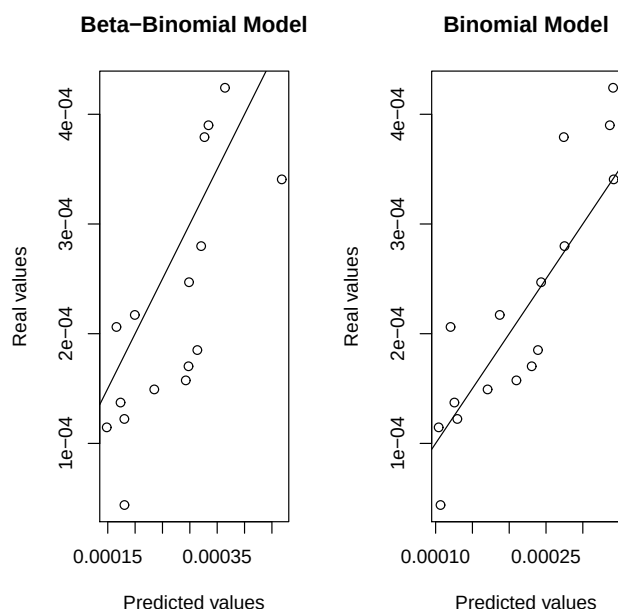
هر دو مدل اتودوجمله‌ای و اتوبتا-دوجمله‌ای با روش INLA و به کمک بسته R-INLA در نرم‌افزار R، بر روی داده‌ها برازش داده شدند. نتایج برازش هر دو مدل در جدول ۲.۵ گزارش شده‌اند. همچنین شکل ۱.۵ نمودارهای پراکنش مقادیر واقعی نسبت بیماران مبتلا به سرطان معده را در مقابل مقادیر برازش‌شده از هر دو مدل، نمایش می‌دهد. با توجه به نتایج جدول و نمودارهای پراکنش مذکور، می‌توان گفت

- هر دو متغیر تبیینی معنی‌دار هستند.
- برآوردهای پارامترها از نظر اندازه اثر به هم نزدیک هستند ولی دقت برآوردها (بر حسب انحراف معیارهای برآوردشده و فاصله‌های اعتبار ۸۰ درصد) اختلاف قابل ملاحظه با هم دارند. با توجه به عدم انعطاف لازم در مساله مدل‌بندی بیش‌پراکنش یا کم‌پراکنش موجود در داده‌ها توسط مدل دوجمله‌ای، انحراف معیارهای کوچک در این مدل می‌تواند نشانه‌ای از کم‌برآورد کردن دقت واقعی برآوردگرها باشد. از این منظر، مدل بتا-دوجمله‌ای توانایی ارائه توصیف بهتری از مدل را دارد. البته مساله نیاز به بررسی بیش‌تر دارد.
- طول مدت اقامت بیماران بر نرخ سرطان معده در شهرستان‌های استان گیلان، تاثیر مستقیم (مثبت) دارد. با توجه به تفسیر پارامترها در مدل‌های لجیت، برآورد ۰/۶۲ برای اثر طول مدت بستری در مدل بتا-دوجمله‌ای، به این معنی است که اگر طول روزهای بستری بیماران ساکن در شهرستانی یک روز افزایش یابد، بخت ابتلا به سرطان معده در آن شهرستان ۸۶ درصد افزایش می‌یابد.
- در مقابل، نرخ سرانه فضای سبز در هر دو مدل تاثیر منفی دارد؛ به این معنی که اگر نرخ سرانه فضای سبز در شهرستانی یک واحد افزایش یابد، بخت ابتلا به سرطان معده در آن شهرستان ۴۶ درصد کاهش خواهد داشت.
- با توجه به نمودارهای پراکنش مقادیر واقعی در مقابل برازش‌شده نسبت ابتلا به سرطان، به سختی می‌توان برازش بهتر نسبی مدل دوجمله‌ای را در مقابل مدل بتا-دوجمله‌ای مشاهده کرد. البته این برازش بهتر می‌تواند منبعی برای بیش‌برازش مدل دوجمله‌ای باشد که در بحث پیش‌گویی فضایی منتج به عملکرد ضعیف خواهد شد. برای بررسی این موضوع مقادیر پاسخ دو شهرستان آستانه اشرفیه و لنگرود را کنار گذاشتیم و با بقیه داده‌ها دو مدل را برازش دادیم. سپس از روی مدل‌های برازش‌شده این دو مقدار را پیش‌گویی کردیم. مقادیر میانگین توان دوم خطای پیش‌گویی MSPE برای دو مدل دوجمله‌ای و بتا-دوجمله‌ای به ترتیب برابر ۰/۰۱۸ و ۰/۰۰۴ به‌دست آمدند. در واقع

مدل بتا-دوجمله‌ای بیش از چهار برابر مدل دوجمله‌ای در پیش‌گویی تواناتر و دقیق‌تر عمل کرده است. همان‌طور که اشاره کردیم، دلیل این نتیجه بیش‌برازش بودن مدل دوجمله‌ای در مقابل مدل بتا-دوجمله‌ای است.

جدول ۲.۵: نتایج برازش دو مدل بتا-دوجمله‌ای و دوجمله‌ای برای داده‌های سرطان معده

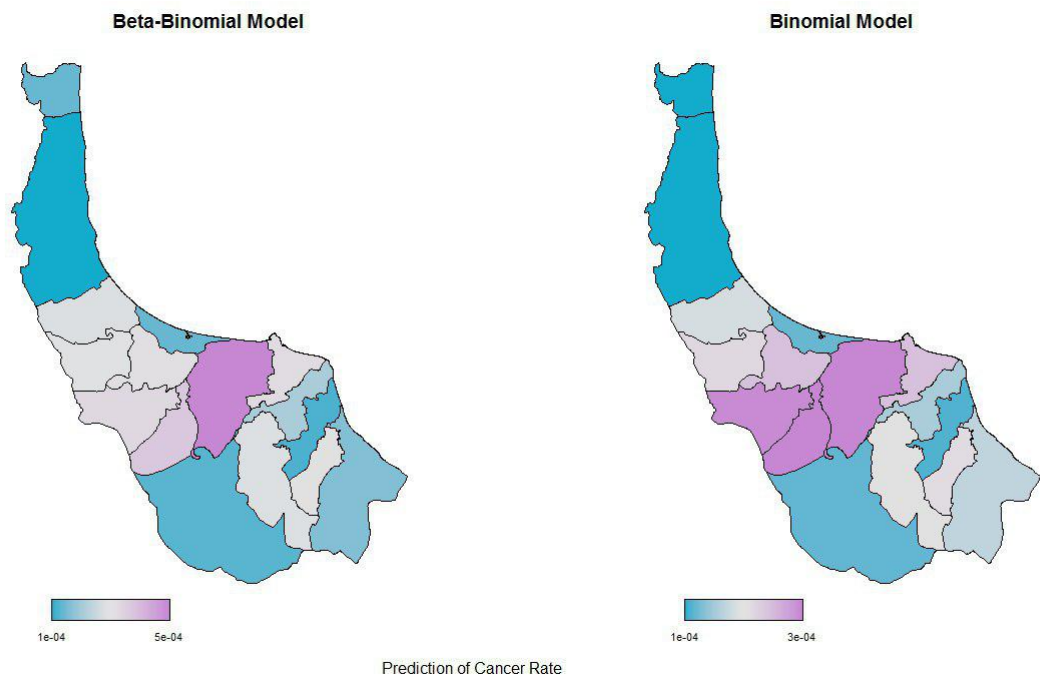
اثر	برآورد اثر	انحراف معیار	ناحیه اعتبار ۸۰٪
ضریب ثابت	-۸/۳۷	۰/۲۳	(-۸/۷۰, -۷/۹۴)
مدل بتا-جمله‌ای			
نرخ مدت اقامت	۰/۶۲	۰/۳۸	(۰/۱۳, ۱/۱۰)
نرخ فضای سبز	-۰/۶۱	۰/۴۰	(-۱/۱۳, -۰/۱۰)
ضریب ثابت	-۸/۵۴	۰/۱۷	(-۸/۶۹, -۸/۴۱)
مدل دوجمله‌ای			
نرخ مدت اقامت	۰/۷۳	۰/۱۸	(۰/۵۲, ۰/۹۵)
نرخ فضای سبز	-۰/۶۹	۰/۱۸	(-۰/۹۲, -۰/۴۶)



شکل ۱.۵: نمودار پراکنش مقادیر برازش‌شده در مقابل مقادیر واقعی نرخ سرطان معده در دو مدل دوجمله‌ای (سمت راست) و بتا-دوجمله‌ای (سمت چپ)

نقشه‌های پهنه‌بندی حاصل از دو مدل برای پیش‌گویی نرخ سرطان معده در استان گیلان در شکل ۲.۵ نشان داده شده‌اند. همان‌طور که از نقشه‌های پهنه‌بندی مشخص است، مقادیر

پیش‌گویی نرخ ابتلا به سرطان در اغلب شهرستان‌های استان گیلان توسط هر دو مدل مشابه هستند؛ در شهرستان‌هایی که اختلاف وجود دارد، مدل دوجمله‌ای نرخ بزرگ‌تری را نسبت به مدل بتا- دوجمله‌ای پیش‌گویی کرده است. ریشه ممکن این مساله را نیز در بیش‌برازش بودن مدل دوجمله‌ای مطرح کردیم. بنابراین، با توجه به توانایی بیش‌تر مدل بتا- دوجمله‌ای در پیش‌گویی، بر استفاده از نتایج مبتنی بر مدل بتا- دوجمله‌ای هم در توصیف عوامل موثر بر افزایش یا کاهش نرخ سرطان و هم پیش‌گویی، تاکید می‌کنیم.



شکل ۲.۵: نقشه‌های پهنه‌بندی نرخ سرطان معده برای مدل‌های دوجمله‌ای (سمت راست) و بتا- دوجمله‌ای (سمت چپ)

۹.۵ نتیجه‌گیری و آینده تحقیق

در این پایان‌نامه برای رفع نقص‌های مدل دوجمله‌ای، به معرفی و بررسی عملکرد مدل کریگینگ بتا- دوجمله‌ای به عنوان جانشینی برای مدل‌بندی پاسخ‌های نرخ همبسته فضایی پرداختیم. با استفاده از یک مطالعه شبیه‌سازی، دقت این مدل در برآورد پارامترها و پیش‌گویی فضایی پاسخ‌های نرخ را نسبت به مدل هم‌کریگینگ دوجمله‌ای، به‌ویژه برای اندازه‌های نمونه نسبتاً کوچک، نشان دادیم. مدل پیشنهادی می‌تواند جانشینی ساده برای مدل‌های پیچیده‌تر فضایی، مانند مدل‌های آمیخته خطی تعمیم‌یافته و مدل‌های سلسله‌مراتبی بیزی برای اندازه نمونه کوچک باشد.

با توجه به مطالبی که در این پایان‌نامه به آن‌ها پرداخته شد، می‌توان به موضوعات زیر

به‌عنوان آینده تحقیق اشاره کرد:

- کریگینگ بتا- دوجمله‌ای می‌تواند برای مدل‌بندی داده‌های دودویی به کار گرفته شود. مدلی که به‌طور معمول برای مدل‌سازی داده‌های دودویی استفاده می‌شود، کریگینگ نشانگر است. در کریگینگ نشانگر فرض می‌شود که متغیر تصادفی دودویی، نتیجه یک متغیر آستانه‌ای است نه یک مشاهده دودویی تصادفی. کریگینگ بتا- دوجمله‌ای می‌تواند برای تشخیص تغییرات اضافی فرآیند دودویی و روند زیربنایی فضایی مفید باشد. همچنین با توجه به عملکرد مناسب مدل بتا- دوجمله‌ای در اندازه نمونه کوچک، ممکن است روش مناسبی برای مدل‌سازی داده‌های دودویی باشد.
- بررسی کریگینگ بیزی مدل بتا- دوجمله‌ای فضایی.
- تعمیم مدل کریگینگ بتا- دوجمله‌ای به مدل کریگینگ بتا- دوجمله‌ای تعمیم‌یافته.

مراجع

- [۱] رضانی ب. و حنیفی ا، (۱۳۸۷) ”شناخت پراکندگی جغرافیایی شیوع سرطان معده در استان گیلان“، **علوم و تکنولوژی محیط زیست**، شماره دو، دوره سیزدهم.
- [۲] میرزاجانی م، (۱۳۹۵)، پایان نامه ارشد: ”استنباط بیزی تقریبی مدل های رگرسیونی پویا با استفاده از تقریب لاپلاس آشیانه ای جمع بسته“، دانشکده علوم ریاضی، دانشگاه صنعتی شاهرود.
- [۳] محمدزاده م، (۱۳۹۱)، ”آمار فضایی و کاربردهای آن“، چاپ اول، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران.
- [۴] حسن زاده ف و کاظمی ا، (۱۳۹۰)، ”مقایسه ی انواع مدل های رگرسیون در تحلیل داده های شمارشی“، **مجله بررسی های رسمی ایران**، شماره ۱، سال ۲۲، ۳۲-۱۷.
- [5] Ansley, C. F. (1979), ”An Algorithm for Exact Likelihood of a Mixed Autoregressive Moving Average Process”, **Biometrika**, 66, 59-65.
- [6] Bartlett, M. S. (1971), ”Physical Nearest-Neighbour Models and Non-Linear Time-Series”, **Journal of Applied Probability**, 8, 222, 232.
- [7] Besag, J. (1974), ”Spatial Interaction and Statistical Analysis of Lattice System (with Discussion)”, **Journal of the Royal Statistical Society**, B, 36(2), 192-236.
- [8] Besag, J. and Moran, P. A. P. (1975), ”On the Estimation and Testing of Spatial Interaction Gaussian Lattice Process”, **Biometrika**, 62, 555-262.
- [9] Billheimer, D., Cardoso, T., Freeman, E., Guttorp, P., Ko, H. W., and Silkey, M. (1997) ”Natural Variability of Benthic Species Composition in the Delaware Bay”, **Environmental and Ecological Statistics**, 4(2), 95-115.
- [10] Branscum, A. J., Johnson W. O. and Thurmond, M. (2007). ”Bayesian Beta Regression: Applications to Household Expenditure Data and Genetic Distance Between Food and Mouth Diseases Viruses”, **Australian and New Zealand Journal of Statistics**, 49, 287-301.

- [11] Breslow, N. E. and Clayton, D. G. (1993), "Approximate Inference in Generalized Linear Mixed Models" **Journal of the American Statistical Association**, 88(421), 9-25.
- [12] Kaiser, M.S. and Cressie, N. (1997), "Modeling Poisson variables with positive spatial dependence", **Statistics and Probability Letters**, 35(4), pp.423-432.
- [13] Carlin, B. P. and Benarjee, S. (2003), "Hierarchical Multivariate CAR Models for Spatio-Temporally Correlated Data", **In Bayesian Statistics 7** (eds. J. M. Bernardo et.al), 45-63, Oxford University Press.
- [14] Chen, F. (2002), "Bayesian Modeling of Multivariate Spatial Binary Data With Application to the Distribution of Plant Species" , PHD Thesis, Department of Statistics, Florida State University.
- [15] Clayton, D. and Kaldor, J. M. (1987), "Empirical Bayes Estimates of Age-Standardized Relative Risk for Use in Disease Mapping" , **Biometrics**, 43, 671-681.
- [16] Noel Cressie. (1985), "Fitting variogram models by weighted least squares", **Journal of the International Association for Mathematical Geology**, 17 (5), 563-586.
- [17] Cressie, N. (1990), "The Origins of Kriging", **Elsevier**, 22, 239-252.
- [18] Cressie, N. (1993), "**Statistics for Spatial Data**", Revised Edition, Wiley, New York.
- [19] Cressie, N. (1988), " Spatial Prediction and Ordinary Kriging", **Mathematical Geology**, 20, 405-421.
- [20] David, M. (1977), " Geostatistical ore Reserve Estimation Elsevier", **Elsevier**, Amsterdam.
- [21] Dent, W. and Min, A. (1978), " A Monte Carlo Study of Autoregressive Moving Average Processes", **Jornal of Econometrie**, 7, 23-55.
- [22] Diggle, P. J., Moyeed, R. A. and Tawn, J. A. (1998), " Model-Based Geostatistics", **Journal of the Royal Statistical Society**, 47(3), 299-350.
- [23] Gandin, L. S. (1963), " Objective Analysis of Meteorological Fields", **Gidrometeor, Leningrad**.
- [24] Gelfand, A. E, and Vounatsou, P. (2003), "Proper Multivariate Conditional Autoregressive Models for Spatial Data Analysis", **Biostatistics**, 1, 11-15.

- [25] Goovaerts, P. (1998), " Ordinary Cokriging Revisited ", **Mathematical Geology**, 30, 21–42, 1998.
- [26] Gotway, A.C ., Carol, A. and Walter, W. S. (1997), "A Generalized Linear Model Approach to Spatial Data Analysis and Prediction" **Journal of Agricultural, Biological, and Environmental Statistics**, 2(2), 157–178.
- [27] Green, J. D. (1970), "Personal Media Probabilities", **Journal of Advertising Research**, 10, 12-18.
- [28] Griffiths, D. A. (1973), "Maximum Likelihood Estimation for the Beta-Binomial Distributions and an Application to the household Distribution of the Total Number of Cases of a Disease", **Biometrics**, 29, 637–648.
- [29] Hainning, R. P. (1978), " The Moving Avrage for Spatial Intraction" **Tranctions, Institute of British Geographers**, N. S. 3, 202-225.
- [30] Huynth, H. (1979), "Statistical Inference for Two Reliability Indices in Mastery Testing Based on the Beta-Binomial Model", **Journal of Educational Statistics**, 4, 231-246.
- [31] Jakson, M. C (2003), "Spatial Data Analysis for Discrete Data on a Lattice", Thesis of PHD, University of Mariland, Department of Mathematics.
- [32] Journel, A. G. and Huijbregts, C. J. (1978), "Mining Geostatistics", **Academic Press**, London.
- [33] Journe, A. G. and Rossi, M. E. (1989), "When do we Need a Trend Model in Kriging?", **Mathematical Geology**, 21, 715–739.
- [34] Kupper, L. L, and Haseman, J. K. (1978), "The use of a correlated binomial model for the analysis of certain toxicological experiments", **Biometrics** , 69-76.
- [35] Kupper, L. L., Portier, C., Hogan, M. D., and Yamamoto, E. (1986), "The Impact of Litter Effects on Dose-Response Modeling in Teratology", **Biometrics**, 42, 85-98.
- [36] Lajaunie, C. (1991), "Local Risk Estimation for a Rare Noncontagious Disease Based on Observed Frequencies", **Note N-36/91/G, Centre de Géostatistique, Ecole des Mines de Paris**.
- [37] Lawson, A. B., Biggeri, A. B., Boehning, D., Lesaffre, E., Viel, J. F., Clark, A., and Divino, F. (2001). " Disease Mapping Models: an Empirical Evaluation" **Statistics in Medicine**, 2217-2241.

-
- [38] Lesauskiene, E. and Ducinskas, k. (2003), " Universal Kriging for Spatio-Temporal Data" **Mathematical Modelling and Analysis**, 8(4), 283–290, 2003.
- [39] Lord, F. M. (1965), "A Strong True-Score Theory, with Applications", **Psychometrika**, 30, 234-270.
- [40] Mardia, K. V. (1988), " Multi-Dimentional Multivariate Gaussian Markov Random Fields with Applications to Image Processing", **Journal of Multivariate Analysis**, 24, 265-284.
- [41] Massy, William F., David B. Montgomery, and Donald G. Morrison. (1970), "Stochastic Models of Buying Behavior (Cam-bridge, Mass)".
- [42] Matheron, G. (1962), "Traite de Geostatistique Appliquee", **Tome I, Memoires du Bureau de Recherches Geologiques et Minieres**, 14, Paris.
- [43] Matheron, G. (1963), "Principles of Geostatistics, with Applications", **Economic Geology**, 58(8), 1246-1266.
- [44] Mazzarino, M. J., Heinrichs, H., and Folster, H. (1983), " Holocene Versus Accelerated Actual Proton Consumption in German Forest Soils", **Effects of Accumulation of Air Pollutants in Forest Ecosystems. Reidel, Dordrecht**, 64, 531-534.
- [45] McLoad, A. I. (1977)," Improved Box-Jenkins Estimators", **Biometrika**, 64, 531-534.
- [46] McCullach, P. and Nelder, J. A. (1989)," Generalized Linear Models", Second Edition, Chapman , Hall, London.
- [47] Melanie, M. W. (2002), " A Close Look at the Spatial Structure Impied by the Car and Sar Models" **Journal of Statistical Planning and Inference**, 121, 311-324.
- [48] Monestiez, P., Dubroca, L., Bonnin, E., Durbec, J. P., and Guinet, C. (2006). "Geostatistical Modelling of Spatial Distribution of Balaenoptera Physalus in the Northwestern Mediterranean Sea from Sparse Count Data and Heterogeneous Observation Efforts", **Ecological Modelling**, 193, 615-628.
- [49] Nelder, J. A. and Wedderburn, R. W. M. (1972), " Generalized Linear Models", **Journal of the Royal Statistical Society, Series A**, 135(3), 370-384.
- [50] Nelder, J . A and Baker, R. J. (1972), " Generalized Linear Models", **In Encyclopedia of Statistical Sciences**, chapter 4.

- [51] Oliver, M.A., Webster, R., Lajaunie, C., Muir, K.R., Parkes, S.E., Cameron, A.H., Stevens, M.C.G. and Mann, J.R.,. (1998), "Binomial Cokriging for Estimating and Mapping the Risk of Childhood Cancer", **Mathematical Medicine and Biology**, 15, 279-297.
- [52] Pack, S. E. (1986), "Hypothesis Testing for Proportions with Overdispersion", **Biometrics**, 42, 967-972.
- [53] Paul, S. R. (1982), "Analysis of Proportions of Affected Fetuses in Teratology Experiments", **Ecological Modelling**, 38, 361-370.
- [54] Pearson, E. S. (1925), " Bayes' theorem, Examined in the Light of Experimental Sampling", **Biometrics**, 17, 388-442.
- [55] Pettitt, A. N, Weir, I. S. and Hart, A. G. (2002), "A Conditional Autoregressive Gaussian Process for Irregularly Spaced Multivariate Data with Application to Modeling Large Sets of Data." **Statistics and Computing**, 4, 95-115.
- [56] Ripley, B. D. (1981), "**Spatial Statistics**", John Wiley, New York.
- [57] 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 B**, 71, 319–392.
- [58] Sain, R. E, and Cressie, N. (2004), "A Spatial Model for Multivariate Lattice Data", Ohio State. University.
- [59] Schuckers, M. E. (2003), " Using the Beta-Binomial Distribution to Assess Performance of a Biometric Identification Device", **Nternational Journal of Image and Graphics**, 3(03), 523-529.
- [60] Skellam, J. G. (1948), " A Probability Distribution Derived from the Binomial Distribution by Regarding the Probability of Success as a Variable Between the Sets of Trials", **Journal of the Royal Statistical Society. Series B (Methodological)**, 10, 257-261.
- [61] Solow, A. R. (1986), " Mapping by Simple Indicator Kriging", **Mathematical Geology**, 18, 335–352.
- [62] Snow, J. (1854), " he cholera near Golden-Square, and at Deptford", **Medical Times and Gazette**, 9, pp.321-322..

- [63] Tarone, R. E. (1982), "The Use of Historical Control Information in Testing for a Trend in Proportions", **Ecological Modelling**, 34, 69-76.
- [64] Tripathi, R. C. Gupta, R. C and Gurland. J. (1994), "Estimation of Parameters in the Beta Binomial Model", **Annals of the Institute of Statistical Mathematics**, 46(2), 317–331.
- [65] Whittle, P. (1954), " On Stationary Processes in the Plance", **Biometrika**, 434-449.
- [66] Wilkinson D, Tanser F. (1999) " GIS/ GPS to Document Increased Access to Community-Based Treatment for Tuberculosis in Africa.", **The Lancet**, 354, 394-395.
- [67] Williams, D. A. (1975), " The Analysis of Binary Responses from Toxicological Experiments Involving Reproductionand Teratogenicity", **Biometrics**, 31, pp. 949–952.
- [68] Yahav, I, Shmueli, G. (2007), "An Elegant Method for Generating Multivariate Poisson Random Variables", **rXiv:0710**, 5670 [stat.CO].
- [69] Zimmerman, D., Pavlik, C, Ruggles, A., and Armstrong, M. P. (1999), " An Experimental Comparison of Ordinary and Universal Kriging and Inverse Distance Weighting", **Mathematical Geology**, 31, 375–390.

پیوست آ

دستورات نرم افزار R

در این پیوست دستورات مربوط به شبیه سازی میدان تصادفی بتا- دوجمله ای، برآورد پارامترها و پیش گویی فضایی شبیه سازی شده در فصل سوم، ارائه شده اند.

```
library(RandomFields)
BETA<-c(1,3)
SILL<-1
RANGE<-10
NUGGET<-0
SAMPLE<-c(2,4,8)
N<-1000
#simulate a Gassian random field
lat<- runif(n=N,min=0,max=20)
lat
lon<-runif(n=N,min=0,max=20)
lon
coord<-cbind(lat,lon)
coord
```

```

colnames(coord)<-c('lat','lon')
colnames(coord)
nobs<-length(lat)
nobs
model<- RMspheric(var=SILL, scale=RANGE)+ RMnugget(var=NUGGET)
sim.normal<- RFsimulate(model=model, x=lat, y=lon, grid=FALSE)
sim.normal
model
mean<-mean(sim.normal$variable1)
mean
sd<-sd(sim.normal$variable1)
sim.cdf<-pnorm(sim.normal$variable1,mean=mean, sd=sd)
sim.beta<-qbeta(sim.cdf,shape1=BETA[1],shape2=BETA[2])
sim.beta<-cbind(coord,sim.beta)
colnames(sim.beta)<-c('lat','lon','B')
sim.beta<- as.data.frame(sim.beta)
#Generate binomial
sim.binomial<- matrix(nrow=nobs, ncol=2)
for(I in 1:nobs){
size<-sample(SAMPLE,1)
sim.binomial[I,1]<-rbinom(n=1,size=size,prob=sim.beta$B[I])
sim.binomial[I,2]<-size
}
sim.binomial<-as.data.frame(cbind(coord,sim.binomial))
colnames(sim.binomial)<-c('lat','lon','Z','N')
sim.binomial$P<-sim.binomial$Z/sim.binomial$N
sim.binomial$P

```

دستورات لازم برای برآورد پارامترهای شکل بتا

```

library(RandomFields)
library(geoR)
library(sp)
BETA<-c(5,1)
SILL<-1

```

```
RANGE<-10
NUGGET<-0
SAMPLE<-2:5
N<-500
est.param<-matrix(nrow=1000,ncol=6)
colnames(est.param)<-c('alpha_small','beta_small','alpha_medium','beta_medium',
'alpha_large','beta_large')
for(RUN in 1:1000){
lat<-runif(n=N,min=0,max=20)
lon<-runif(n=N,min=0,max=20)
coord<-cbind(lat,lon)
nobs<-N
colnames(coord)<-c('lat','lon')
model<-RMspheric(var=SILL,scale=RANGE)+RMnugget(var=NUGGET)
sim.normal<-RFsimulate(model=model,x=lat,y=lon,grid=FALSE)
mean<-mean(sim.normal$variable1)
sd<-sd(sim.normal$variable1)
sim.cdf<-pnorm(sim.normal$variable1,mean=mean,sd=sd)
sim.beta<-qbeta(sim.cdf,shape1=BETA[1],shape2=BETA[2])
sim.beta<-cbind(coord,sim.beta)
colnames(sim.beta)<-c('lat','lon','B')
sim.beta<-as.data.frame(sim.beta)
sim.binary.small<-matrix(nrow=nobs,ncol=2)
small<-2:5
sim.binary.medium<-matrix(nrow=nobs,ncol=2)
medium<-10:20
sim.binary.large<-matrix(nrow=nobs,ncol=2)
large<-50:100
for(I in 1:nobs){
size.s<-sample(small,1)
sim.binary.small[I,1]<-rbinom(n=1,size=size.s,prob=sim.beta$B[I])
sim.binary.small[I,2]<-size.s
}
lat<-runif(n=N,min=0,max=20)
```

```

lon<-runif(n=N,min=0,max=20)
coord<-cbind(lat,lon)
nobs<-N
colnames(coord)<-c('lat','lon')
model<-RMspheric(var=SILL,scale=RANGE)+RMnugget(var=NUGGET)
sim.normal<-RFsimulate(model=model,x=lat,y=lon,grid=FALSE)
mean<-mean(sim.normal$variable1)
sd<-sd(sim.normal$variable1)
sim.cdf<-pnorm(sim.normal$variable1,mean=mean,sd=sd)
sim.beta<-qbeta(sim.cdf,shape1=BETA[1],shape2=BETA[2])
sim.beta<-cbind(coord,sim.beta)
colnames(sim.beta)<-c('lat','lon','B')
sim.beta<-as.data.frame(sim.beta)
for(I in 1:nobs){
size.m<-sample(medium,1)
sim.binary.medium[I,1]<-rbinom(n=1,size=size.m,prob=sim.beta$B[I])
sim.binary.medium[I,2]<-size.m
}
lat<-runif(n=N,min=0,max=20)
lon<-runif(n=N,min=0,max=20)
coord<-cbind(lat,lon)
nobs<-N
colnames(coord)<-c('lat','lon')
model<-RMspheric(var=SILL,scale=RANGE)+RMnugget(var=NUGGET)
sim.normal<-RFsimulate(model=model,x=lat,y=lon,grid=FALSE)
mean<-mean(sim.normal$variable1)
sd<-sd(sim.normal$variable1)
sim.cdf<-pnorm(sim.normal$variable1,mean=mean,sd=sd)
sim.beta<-qbeta(sim.cdf,shape1=BETA[1],shape2=BETA[2])
sim.beta<-cbind(coord,sim.beta)
colnames(sim.beta)<-c('lat','lon','B')
sim.beta<-as.data.frame(sim.beta)
for(I in 1:nobs){
size.1<-sample(large,1)

```



```
sim.binary.large[I,1]<-rbinom(n=1,size=size.1,prob=sim.beta$B[I])
sim.binary.large[I,2]<-size.1
}
sim.binary.small<-as.data.frame(cbind(coord,sim.binary.small))
colnames(sim.binary.small)<-c('lat','lon','Z','N')
sim.binary.medium<-as.data.frame(cbind(coord,sim.binary.medium))
colnames(sim.binary.medium)<-c('lat','lon','Z','N')
sim.binary.large<-as.data.frame(cbind(coord,sim.binary.large))
colnames(sim.binary.large)<-c('lat','lon','Z','N')
# Estimate beta parameters
# Depending on computer platform (Mac or PC) , family name
# is 'betabinomial .ab' or 'betabinomialff '
library(VGAM)
fit.small<-vglm(cbind(Z,N-Z)~1,betabinomialff,data=sim.binary.small)
est.param[RUN,1]<-Coef(fit.small)[1]
est.param[RUN,2]<-Coef(fit.small)[2]
fit.medium<-vglm(cbind(Z,N-Z)~1,betabinomialff,data=sim.binary.medium)
est.param[RUN,3]<-Coef(fit.medium)[1]
est.param[RUN,4]<-Coef(fit.medium)[2]
fit.large<-vglm(cbind(Z,N-Z)~1,betabinomialff,data=sim.binary.large)
est.param[RUN,5]<-Coef(fit.large)[1]
est.param[RUN,6]<-Coef(fit.large)[2]
}
```

دستورات لازم برای برآورد پارامترهای ساختار وابستگی فضایی

```
rm(list=ls())
library(RandomFields)
library(RColorBrewer)
library(geoR)
library(sp)
library(VGAM)
set.seed(2387)
STARTING <- read.csv('starting_values.csv')
N <- 100
```

```

NSIM <- 100
for(s in 1:nrow(STARTING)){
  BETA <- STARTING[s,1:2]
  SILL<-1
  RANGE <- STARTING[s,3]
  NUGGET<-0
  SAMPLESIZE<-STARTING[s,4]:STARTING[s,5]

  est.param<-matrix(nrow=NSIM,ncol=14)
  colnames(est.param)<-c('beta.tau','beta.sigma','beta.phi','bbk1.tau','bbk1.sigma',
    , 'bbk1.phi','bbk2.tau','bbk2.sigma','bbk2.phi','bbk3.tau','bbk3.sigma',
    'bbk3.phi','alpha','beta')
  for(RUN in 1:NSIM){
    print(RUN)
    #Simulate a dense Gaussian random fields
    lat<-seq(from=0,to=20,length=100)
    lon<-seq(from=0,to=20,length=100)
    coord<-expand.grid(lat,lon)
    nobs<-N
    colnames(coord)<-c('lat','lon')
    model<-RMspheric(var=SILL,scale=RANGE)+RMnugget(var=NUGGET)
    sim.normal<-RFsimulate(model=model,x=lat,y=lon,grid=TRUE)
    sim.normal<-cbind(coord,sim.normal$variable1)
    colnames(sim.normal)<-c('lat','lon','A')
    at<-seq(from=-4,to=4,length=100)
    col<-colorRampPalette(brewer.pal(9,'Reds'))(100)
    coordinates(sim.normal)<-~lat+lon
    mean<-mean(sim.normal$A)
    sd<-sd(sim.normal$A)
    sim.cdf<-pnorm(sim.normal$A,mean=mean,sd=sd)
    sim.beta<-qbeta(sim.cdf,shape1=BETA$alpha,shape2=BETA$beta)
    sim.beta<-cbind(coord,sim.beta)
    colnames(sim.beta)<-c('lat','lon','B')
    sim.beta<-as.data.frame(sim.beta)
  }
}

```

```
at<-seq(from=0,to=1,length=100)
coordinates(sim.beta)<-~lat+lon
gridded(sim.beta)<-TRUE
spplot(sim.beta,zcol='B',main=paste('Simulated beta distribution,alpha=',BETA[1],
'beta=',BETA[2]),at=at,col.regions=col,colorkey=TRUE)
sample.points<-sample(x=1:nrow(coord),size=N)
beta.points<-sim.beta[sample.points,]
coord.points<-coord[sample.points,]
sim.binomial<-matrix(nrow=N,ncol=2)
for(I in 1:N){
  size<-sample(SAMPLESIZE,1)
  sim.binomial[I,1]<-rbinom(n=1,size=size,prob=beta.points$B[I])
  sim.binomial[I,2]<-size
}
sim.binomial<-as.data.frame(cbind(coord.points,sim.binomial))
colnames(sim.binomial)<-c('lat2','lon2','Z','N')
sim.binomial$P<-sim.binomial$Z/sim.binomial$N
fit.ab<-vglm(cbind(Z,N-Z)~1,family=betabinomialff,data=sim.binomial)
est.param[RUN,13]<-Coef(fit.ab)[1]
est.param[RUN,14]<-Coef(fit.ab)[2]
alpha<-Coef(fit.ab)[1]
beta<-Coef(fit.ab)[2]
bias<-alpha*beta/((alpha+beta)*(alpha+beta+1))
#Estimate beta-binomial kriging variogram
varS<-matrix(nrow=N*N,ncol=3)
colnames(varS)<-c('dist','value','mult')
index<-1
for(i in 1:N){
  for(j in 1:N){
    mult<-(sim.binomial$N[i]*sim.binomial$N[j])/(sim.binomial$N[i]+sim.binomial$N[j])
    varS[index,1]<-dist(sim.binomial[c(i,j),1:2])
    varS[index,2]<-mult*(dist(sim.binomial$P[c(i,j)])^2)-bias
    varS[index,3]<-mult
    index<-index+1
  }
}
```

```

}
}
varS<-as.data.frame(varS)
bin.varP<-variog(data=sim.binomial$P,coords=coord.points,breaks=1:20)
beta.var<-variog(data=sim.beta$B,coords=coord,breaks=1:20)
#Use lag distances as defined in 'variog'
h<-1:20
bin.varS<-matrix(nrow=length(h),ncol=4)
colnames(bin.varS)<-c('nobs','var','h','ave.weights')
for(H in h){
x<-varS[varS$dist<H& varS$dist>=(H-1),]
bin.varS[H,1]<-nrow(x)
bin.varS[H,2]<-max(colSums(x)[2]/(2*colSums(x)[3]),0)
bin.varS[H,3]<-H
bin.varS[H,4]<-mean(1/mult)
}
vario.beta<-beta.var
vario.prop<-bin.varP
vario.bbk<-bin.varP
vario.bbk$v<-bin.varS[,2]
vario.bbk2<-bin.varP
vario.bbk2$v<-bin.varS[,2]
vario.bbk2$n<-bin.varS[,4]*bin.varS[,1]
beta.parm <- variofit(vario=vario.beta, ini.cov.pars=c(0.1,RANGE), cov.model='sph',
fix.nugget=FALSE)
bbk.parm1 <- variofit(vario=vario.bbk, ini.cov.pars=c(0.1,RANGE), cov.model='sph',
fix.nugget=FALSE,weights='npairs')
bbk.parm2 <- variofit(vario=vario.bbk2, ini.cov.pars=c(0.1,RANGE), cov.model='sph',
fix.nugget=FALSE,weights='npairs')
bbk.parm3 <- variofit(vario=vario.bbk, ini.cov.pars=c(0.1,RANGE), cov.model='sph',
fix.nugget=FALSE,weights='cressie')
est.param[RUN,1]<-beta.parm$nugget
est.param[RUN,2]<-beta.parm$cov.pars[1]
est.param[RUN,3]<-beta.parm$cov.pars[2]

```

```
est.param[RUN,4]<-bbk.parm1$nugget
est.param[RUN,5]<-bbk.parm1$cov.pars[1]
est.param[RUN,6]<-bbk.parm1$cov.pars[2]
est.param[RUN,7]<-bbk.parm2$nugget
est.param[RUN,8]<-bbk.parm2$cov.pars[1]
est.param[RUN,9]<-bbk.parm2$cov.pars[2]
est.param[RUN,10]<-bbk.parm3$nugget
est.param[RUN,11]<-bbk.parm3$cov.pars[1]
est.param[RUN,12]<-bbk.parm3$cov.pars[2]
write.csv(est.param,paste('param',s,'.csv'))
}
}
```

دستورات لازم برای پیش‌گویی فضایی

```
library(RandomFields)
library(geoR)
library(sp)
library(VGAM)

# sim.binomial is the simulated binomial sample data
colnames(sim.binomial) <- c('lat','lon', 'Z','N','P')
fit.ab <- vglm(cbind(Z,N-Z)~1,family=betabinomialfff,data=sim.binomial)
alpha <- Coef(fit.ab)[1]
beta <- Coef(fit.ab)[2]
bias <- alpha*beta/((alpha+beta)*(alpha+beta+1))
# Estimate beta?binomial variogram
varS <- matrix(nrow=N*N,ncol=3)
colnames(varS) <- c('dist','value','mult')
index<-1
for(i in 1:N){
  for(j in 1:N){
    mult <- (sim.binomial$N[i]*sim.binomial$N[j])/
      (sim.binomial$N[i]+sim.binomial$N[j])
    varS[index,1] <- dist(sim.binomial[c(i,j),1:2])
```

```

varS[index,2] <- mult*(dist(sim.binomial$P[c(i,j)])^2)-bias
varS[index,3] <- mult
index <- index+1
  }
}
varS <- as.data.frame(varS)

bin.varP <- variog(data=sim.binomial$P,coords=coord.points,breaks=1:20)
beta.var <- variog(data=sim.beta$B,coords=coord,breaks=1:20)
# Use lag distances as defined in 'variog '
h <- 1:20
bin.varS <- matrix(nrow=length(h),ncol=4)
colnames(bin.varS) <- c('nobs','var','h','ave.weights')
for(H in h){
  x<-varS[varS$dist<H & varS$dist>=(H-1),]
  bin.varS[H,1]<-nrow(x)
  bin.varS[H,2]<-max(colSums(x)[2]/(2*colSums(x)[3]),0)
  bin.varS[H,3]<-H
  bin.varS[H,4]<-mean(1/mult)
}

vario.beta <- beta.var
vario.prop <- bin.varP
vario.bbk2 <- bin.varP
vario.bbk2$v <- bin.varS[,2]
vario.bbk2$n <- bin.varS[,4]*bin.varS[,1]
beta.parm <- variofit(vario=vario.beta,ini.cov.pars=c(0.08,RANGE),
cov.model='sph',fix.nugget=FALSE)
bbk.parm <- variofit(vario=vario.bbk2,ini.cov.pars=c(0.08,RANGE),
cov.model='sph',fix.nugget=FALSE,weights='npairs')
if(bbk.parm$nugget<0) bbk.parm$nugget=0
##MAKE predictions
A <- diag(c(bias/sim.binomial$N,0))
B <- matrix(rep(1,(N+1)^2),nrow=(N+1),ncol=(N+1))

```

```
for(i in 1:N){
  for(j in 1:N){
    dist <- dist(sim.binomial[c(i,j),1:2])
    B[i,j] <- bbk.parm$cov.pars[1]*(1-(3*dist/(2*bbk.parm$cov.pars[2])-
    dist^3/(2*bbk.parm$cov.pars[2]^3)))
    if(dist>bbk.parm$cov.pars[2]) B[i,j] <- 0
  }
}
B[N+1,N+1] <- 0
C <- A+B
Cinv <- solve(C)
# Set up C0 matrix for each point at location o,
# solve for the weights , and predict
pred <- matrix(nrow=nrow(coord),ncol=4)
colnames(pred) <- c('lat','lon','pred','var')
for(o in 1:nrow(coord)){
  C0<-matrix(rep(1,N+1),ncol=1)
  for(i in 1:N){
    dist <- dist(rbind(sim.binomial[i,1:2],coord[o,1:2]))
    C0[i,1] <- bbk.parm$cov.pars[1]*(1-(3*dist/(2*bbk.parm$cov.pars[2])-
    dist^3/(2*bbk.parm$cov.pars[2]^3)))
    if(dist>bbk.parm$cov.pars[2]) C0[i] <- 0
  }
  lambda <- Cinv%*%C0
  pred[o,3] <- max(min(sum(lambda[1:N]*sim.binomial$P),1),0)
  pred[o,1] <- coord[o,1]
  pred[o,2] <- coord[o,2]
  pred[o,4] <- bbk.parm$cov.pars[1]-sum(lambda*C0)
}

pred <- as.data.frame(pred)
pred$beta <- sim.beta$B
pred$error <- pred$pred-sim.beta$B
```

Abstract

The popular model for spatial count responses in finite populations is constructed based on binomial distribution. In many applications, due to overdispersion problem, to use binomial distribution is not appropriate and can result in inefficient inferences regarding the estimation of parameters and spatial prediction. In order to eliminate the defects of a binomial model, an alternative approach for modeling spatially correlated proportions is to use a beta-binomial distribution. The flexibility of beta-binomial distribution could make it the first choice for modeling overdispersed count responses in finite populations. In this thesis, we describe the spatial prediction of rates by using the flexible beta-binomial model and generalize it for spatial analysis of lattice data in the framework of Bayesian inference. We also demonstrate its efficacy by a simulation study. To evaluate the performance of this model, we analyze two real data sets including the number of rainfall days in Semnan province, Iran and the number of people with gastric cancer in Guilan province, Iran.

Keywords: Spatial counts, beta-binomial process, overdispersion, spatial prediction.



Shahrood University of Technology

Faculty of Mathematical Sciences

MSc Thesis in: Mathematical Statistics

A flexible model for spatial prediction of rate responses by using beta-binomial distribution

By: Leila Abedinpour Liarjadmeh

Supervisor

Dr. Hossein Baghishani

Advisor

Dr. Negar Eghbal

January 2018