

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ



دانشکده علوم ریاضی
پایان نامه کارشناسی ارشد آمار ریاضی

مدل بندی سه بعدی ژئومکانیکی بیزی با استفاده از مدل های زمین آماری: مطالعه موردی یکی از مخازن نفتی ایران

نگارنده: فتانه فخری

استاد راهنما
دکتر حسین باغیشنی

استاد مشاور
دکتر بهزاد تخمچی

مهر ۱۳۹۹



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 فتانه فخری با شماره دانشجویی ۹۶۳۵۰۴۴ رشته آمار ریاضی گرایش ----- تحت عنوان مدل بندی سه بعدی ژئومکانیکی بیزی با استفاده از مدل‌های زمین آماری مطالعه موردی یکی از مخازن نفتی ایران که در تاریخ ۱۳۹۹/۰۷/۲۹ با حضور هیأت محترم داوران در دانشگاه صنعتی شاهرود برگزار گردید به شرح ذیل اعلام می‌گردد:

الف) درجه عالی: نمره ۱۹-۲۰ ☒ ب) درجه خیلی خوب: نمره ۱۸-۱۹ ☐ ج) درجه خوب: نمره ۱۷-۱۶ ☐ د) درجه متوسط: نمره ۱۵-۱۴ ☐ ه) کمتر از ۱۴ غیر قابل قبول و نیاز به دفاع مجدد دارد ☐ نوع تحقیق: نظری ☒ عملی ☐			
عضو هیأت داوران	نام و نام خانوادگی	مرتبه علمی	امضاء
۱- استاد راهنمای اول	دکتر حسین باغیشتی	استادیار	
۲- استاد راهنمای دوم	-----	-----	-----
۳- استاد مشاور	دکتر بهزاد تخم چی	دانشیار	
۴- نماینده تحصیلات تکمیلی	دکتر محمدهادی توری اسکندری	استادیار	
۵- استاد ممتحن اول	دکتر محمدرضا ربیعی	استادیار	
۶- استاد ممتحن دوم	دکتر داود شاهسونی	دانشیار	

نام و نام خانوادگی رئیس دانشکده: دکتر مهدی قونمند

تاریخ و امضاء و مهر دانشکده:



با تمام کاستی‌ها تقدیم به
خدای مهربان
و خانواده‌ی عزیزم

سپاس‌گزاری

سپاس بی‌کران پروردگار یکتا را که هستی‌مان بخشید، به طریق علم و دانش رهنمون‌مان شد، به هم‌نشینی رهروان دانش مفتخرمان نمود و خوشه‌چینی از علم و معرفت را روزی‌مان ساخت.

با امتنان بیکران از مساعدت‌های بی‌شائبه‌ی اساتید بزرگوار خود، دکتر حسین باغیشنی و دکتر بهزاد تخم‌چی صمیمانه کمال تشکر را دارم، که نه تنها در نوشتن این پایان‌نامه و فراگیری علم با سعه‌صدر مرا یاری دادند و از هیچ کمکی دریغ ننمودند، بلکه درس اخلاق، فروتنی، استقامت و امید را نیز به من دادند.

از استاد بزرگوار جناب آقای دکتر داوود شاهسونی که یک ترم سرکلاس ایشان حضور داشتم و در فراگیری علوم جدید آماری به بنده بسیار کمک کردند و زحمت داوری این پایان‌نامه را نیز بر عهده گرفتند صمیمانه تشکر می‌کنم.

از استاد محترم، جناب آقای دکتر محمدرضا ربیعی که لطف کردند و داوری این پایان‌نامه را بر عهده گرفتند صمیمانه تشکر مینمایم.

و در نهایت از خانواده عزیزم و دوستانی که همیشه کنارم بوده‌اند، سپاسگزارم. از خداوند مهربانم برای شما عزیزان، آرزوی سلامتی و توفیق روز افزون را مسألت دارم.

فتانه فخری

مهر ۱۳۹۹

تعهد نامه

اینجانب **فتانه فخری** دانشجوی کارشناسی ارشد رشته آمار ریاضی دانشکده علوم ریاضی دانشگاه صنعتی شاهرود، نویسنده پایان نامه با عنوان **مدل بندی سه بعدی ژئومکانیکی بیزی با استفاده از مدل های زمین آماری: مطالعه موردی یکی از مخازن نفتی ایران**، تحت راهنمایی دکتر **حسین باغیشنی** متعهد می شوم:

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

فتانه فخری

مهر ۱۳۹۹

مالکیت نتایج و حق نشر

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

چکیده

استفاده از ویژگی‌های مکانیکی سنگ برای مطالعه فضاهای زیرزمینی، ژئومکانیک نامیده می‌شود. ژئومکانیک یکی از شاخه‌های علمی مهم در حوزه اکتشاف و حفاری مخازن زیرزمینی، به‌ویژه مخازن نفتی، است. اغلب داده‌های حاصل در مطالعات ژئومکانیکی، از نوع زمین‌آماري هستند که مدل‌بندی و تحلیل آن‌ها نیازمند علم آمار فضایی است. از طرفی، بسیاری از داده‌های زمین‌آماري در مطالعات ژئومکانیکی سه‌بعدی و حجیم هستند که توسعه مدل‌های آمار فضایی و روش‌های محاسباتی کارا برای تحلیل آن‌ها را با چالش مواجه می‌کند. با نگاهی به این واقعیت و برای مقابله با دو چالش ذکرشده، در این پایان‌نامه، برای اولین بار، دو روش پیشنهادی جدید را در حوزه مدل‌بندی سه‌بعدی فضایی ارائه می‌کنیم. اول، یک روش ناپارامتری به نام میانه‌پالایش سه‌بعدی را برای برآورد تابع روند یک میدان فضایی سه‌بعدی معرفی می‌کنیم. با استفاده از این روش، پیشگویی فضایی کریگینگ میانه‌پالایش را نیز توسعه می‌دهیم. همچنین کارایی کریگینگ میانه‌پالایش پیشنهادی را با یک کریگینگ عام که تابع روند در آن با روش اسپلاین برآورد می‌شود، در مثال‌های شبیه‌سازی و واقعی مربوط به داده‌های مقاومت فشاری تک‌محوری سنگ در یکی از مخازن نفتی ایران، مقایسه می‌کنیم. دوم، با استفاده از یک رویکرد بیزی، روش تقریبی بیزی لاپلاس آشیانه‌ای جمع‌بسته (INLA) و ترکیب آن با روش معادلات دیفرانسیل جزئی تصادفی (SPDE) را به حالت سه‌بعدی توسعه می‌دهیم.

کلمات کلیدی: زمین‌آمار، استنباط بیزی، کریگینگ میانه‌پالایش سه‌بعدی، تقریب لاپلاس آشیانه‌ای جمع‌بسته، رهیافت INLA+SPDE.

پیش‌گفتار

از گذشته تا به امروز یکی از مسائل اساسی که افکار بسیاری از پژوهشگران را به خود معطوف کرده است، کشف رفتار مواد و موضوع مورد بررسی در مکان‌های فاقد نمونه است. در واقع، محققان ترجیح می‌دهند قبل از نمونه‌برداری، ذهنیتی نسبی از رفتار مواد در آن مکان داشته باشند. این ترجیح به‌ویژه در علوم زمین که نمونه‌برداری هزینه‌های گزافی را بر عهده‌ی پژوهشگر متحمل می‌سازد، بارزتر است. روشن است که آگاهی بر رفتار مواد در ناحیه تحت مطالعه، می‌تواند موجب کاهش هزینه‌های اقتصادی و صرفه‌جویی در زمان شود. بنابراین، پیشگویی نقش بسزایی در علوم مختلف دارد. در علوم زمین و نفت اطلاعات مورد نیاز برای مدل‌بندی و پیشگویی رفتار سنگ، در یک مکان از پیش تعریف‌شده، با نمونه‌برداری از زمین به‌صورت مستقیم و در مواردی با استفاده از روابطی معلوم، اطلاعات مورد نیاز محاسبه می‌شوند. این نوع از داده‌ها که از لحاظ مفهومی پیوسته هستند، مثالی از داده‌های زمین‌آماری است. برای پیشگویی داده‌های زمین‌آماری باید روند داده‌ها را تشخیص داد. اگر نتوان ضابطه‌ای پارامتری برای تابع روند تشخیص داد، از روش‌های ناپارامتری برای مدل‌بندی آن استفاده می‌شود. در تحقیقات متعددی این نوع از مدل‌سازی‌ها به‌خوبی توسعه داده شده‌اند، اما حجم بالا و سه بعدی بودن داده‌ها در مطالعات ژئومکانیکی، پژوهشگران این حیطه را با مشکلات زیادی روبرو می‌سازد.

یکی از روش‌های ناپارامتری در مدل‌بندی تابع روند که کمتر به آن توجه شده است، روش میانه‌پالایش است که برای داده‌های سه بعدی با حجم بالا از پیچیدگی محاسباتی در اجرا برخوردار است. از طرفی، مشابه همین مشکل را می‌توان در تحلیل بیزی این نوع از داده‌ها دریافت. روش‌های مبتنی بر نمونه‌گیری مانند الگوریتم‌های MCMC روش معمول برای محاسبه‌ی توزیع توأم پسین مدل‌های آماری است. اما این روش‌ها از محدودیت‌های محاسباتی سنگین و مشکلات همگرایی برخوردار است. به همین دلیل، رو و همکارانش در سال ۲۰۰۹ یک روش تقریبی معروف به روش تقریب لاپلاس‌آشیانی جمع‌بسته (INLA) را ارائه کردند. در سال ۲۰۱۵ لیندگرین و همکارانش برای مدل‌بندی کارای داده‌های زمین‌آماری با استفاده از INLA، روش معادلات دیفرانسیل جزئی تصادفی (SPDE) را پیشنهاد دادند. تا قبل از این پژوهش به‌ندرت از رویکرد بیزی با استفاده از تقریب INLA برای داده‌های سه بعدی استفاده شده بود. چالش‌های این مسائل موجب شد که نگارنده را برای بررسی‌های بیشتر در این زمینه وادار سازد. بر اساس مقدمات گفته‌شده، ساختار این پایان‌نامه به‌صورت زیر است:

- فصل اول مروری بر تعاریف و مقدمات آمار فضایی است.

- در فصل دوم مدل‌های سه بعدی زمین‌آماری را معرفی می‌کنیم و ابزار برآورد ساختار وابستگی داده‌ها در یک میدان سه بعدی را تشریح می‌کنیم. همچنین روش کریگینگ معمولی سه بعدی را بر روی داده‌های مقاومت فشاری تک‌محوری سنگ در یکی از مخازن نفتی پیاده می‌کنیم.

- در فصل سوم روش میانه‌پالایش سه بعدی و الگوریتم اجرای آن را معرفی می‌کنیم. سپس با استفاده از مثال‌های شبیه‌سازی و داده‌های مقاومت فشاری تک‌محوری سازند، عملکرد آن را با روش ناپارامتری اسپلاین در پیشگویی فضایی به همراه کریگینگ عام مقایسه می‌کنیم.
- در فصل چهارم، جزییات روش بیزی تقریبی INLA را معرفی می‌کنیم و با استفاده از یک مثال شبیه‌سازی روش $INLA + SPDE$ سه بعدی پیشنهادی را اجرا و عملکرد آن را ارزیابی می‌کنیم. همه کدهای کامپیوتری برای اجرای روش‌های این پایان‌نامه در محیط نرم‌افزار R نوشته شده‌اند.

فهرست مقاله‌های مستخرج از پایان‌نامه

۱. فخری، ف.، باغیشنی، ح.، تخم‌چی، ب. (۱۳۹۷)، ناهمسانگردی هندسی در پیشگویی فضایی، اولین کنفرانس دانشجویی آمار، دانشگاه تربیت مدرس تهران، ایران.
۲. فخری، ف.، باغیشنی، ح.، تخم‌چی، ب. (۱۳۹۸)، پیشگویی فضایی سه بعدی به منظور برآورد مقاومت فشاری تک محوری در یکی از مخازن نفتی ایران، سومین کنفرانس آمار فضایی، زنجان، ایران، ۱۳۲ - ۱۲۴.

فهرست مطالب

۱	آمار فضایی	۱
۲	۱.۱ انواع داده‌های فضایی	۱.۱
۴	۲.۱ مدل‌بندی آماری فضایی	۲.۱
۵	۳.۱ ساختار همبستگی فضایی	۳.۱
۶	۱.۳.۱ تغییرنگار	۱.۳.۱
۸	۲.۳.۱ هم‌تغییرنگار	۲.۳.۱
۹	۳.۳.۱ همبستگی‌نگار	۳.۳.۱
۹	۴.۳.۱ ناهمسان‌گردی	۴.۳.۱
۱۰	۵.۳.۱ خانواده‌های توابع همبستگی فضایی	۵.۳.۱
۱۳	۴.۱ استنباط آماری	۴.۱
۱۳	۱.۴.۱ استنباط مبتنی بر درست‌نمایی	۱.۴.۱
۱۴	۲.۴.۱ استنباط بیزی	۲.۴.۱
۱۶	۵.۱ پیشگویی فضایی	۵.۱
۱۷	۱.۵.۱ کریگینگ معمولی	۱.۵.۱
۱۸	۲.۵.۱ کریگینگ عام	۲.۵.۱
۱۹	۲ مدل‌های سه بعدی زمین‌آماري	۲
۱۹	۱.۲ مقدمه	۱.۲
۲۰	۲.۲ تشریح مسأله	۲.۲
۲۲	۳.۲ ساختار همبستگی فضایی	۳.۲
۲۳	۱.۳.۲ برآورد تغییرنگار مخزن نفتی	۱.۳.۲
۲۵	۲.۳.۲ برآوردگر استوار	۲.۳.۲
۲۶	۳.۳.۲ بررسی ناهمسان‌گردی میدان نفتی	۳.۳.۲
۲۷	۴.۲ پیشگویی سه بعدی مقاومت فشاری تک محوری سازند	۴.۲

۳۱	۳	روش‌های ناپارامتری فضایی در میدان‌های تصادفی سه بعدی
۳۱	۱.۳	میان‌پالایش
۳۲	۱.۱.۳	الگوریتم میان‌پالایش سه بعدی
۳۷	۲.۳	کریگینگ سه بعدی میان‌پالایش
۳۸	۱.۲.۳	تجزیه فرآیند فضایی
۴۰	۲.۲.۳	درون‌یابی و برون‌یابی اثرات میان‌پالایش
۴۲	۳.۲.۳	کریگینگ معمولی روی باقی‌مانده‌ها
۴۳	۳.۳	روش اسپلین
۴۶	۴.۳	مطالعه شبیه‌سازی
۴۶	۱.۴.۳	روند خطی
۵۰	۲.۴.۳	روند غیر خطی
۵۲	۵.۳	مطالعه پیشگویی سه بعدی ناپارامتری مقاومت فشاری تک محوری سنگ
۵۷	۴	رهیافت بیزی تقریبی برای مدل‌بندی فضایی سه بعدی
۵۷	۱.۴	مقدمه
۵۹	۲.۴	مدل‌های گوسی پنهان
۶۱	۳.۴	تقریب لاپلاس آشیانه‌ای جمع بسته
۶۲	۱.۳.۴	تقریب چگالی پسین توأم بردار ابرپارامترها
۶۳	۲.۳.۴	روش توری
۶۴	۳.۳.۴	روش طرح مرکب مرکزی
۷۱	۴.۳.۴	تقریب چگالی پسین کناری عناصر میدان پنهان
۷۵	۴.۴	ملاک‌های ارزیابی
۷۶	۱.۴.۴	ملاک اطلاع انحراف
۷۷	۲.۴.۴	ملاک پیش‌گویی شرطی مولفه‌ها
۷۸	۵.۴	داده‌های فضایی سه بعدی
۷۹	۶.۴	رهیافت معادلات دیفرانسیل جزئی تصادفی
۸۱	۷.۴	رهیافت INLA+SPDE برای مدل‌بندی فضایی سه بعدی
۸۲	۱.۷.۴	مثلت‌بندی ناحیه سه بعدی
۸۴	۸.۴	نتیجه‌گیری

فهرست تصاویر

۷	۱.۱ اجزاء تشکیل دهنده‌ی تغییرنگار
۲.۱	تابع همبستگی نمایی به ازاء $\phi = 1$ (خط ممتد) و $\phi = 1$ (خط چین). محل تقاطع خط افقی باریک ($\rho(h) = 0.05$) با منحنی مقدار دامنه‌ی عملی را نشان می‌دهد.
۱۱	۳.۱ خانواده همبستگی مترن. برای $k = (0.5, 1.5, 2.5)$ (خط ممتد، خط چین و نقطه‌چین به ترتیب) مقدار ϕ طوری انتخاب شده که در هر سه مورد مقدار دامنه‌ی عملی برابر با ۳ شود و این مقدار از تقاطع خط نازک افقی ($\rho = 0.05$) با منحنی به‌دست می‌آید.
۱۲	۴.۱ تابع همبستگی کروی (خط ممتد) و مترن (خط چین) برای مقدار $k = 1.5$. در هر دو مورد مقدار ϕ طوری انتخاب شده است که دامنه‌ی عملی برابر با ۳ شود.
۲۲	۱.۲ نمودار سه بعدی چاه‌ها
۲.۲	محل قرارگیری چاه‌های مجموعه آموزشی (توخالی) و مجموعه آزمون (ستاره‌ای)
۲۳	بر حسب طول و عرض جغرافیایی و محل قرارگیری آن‌ها
۲۴	۳.۲ تغییرنگار یک بعدی
۲۵	۴.۲ اجزاء تشکیل دهنده‌ی تغییرنگار
۵.۲	نمودار سه بعدی نیم‌تغییرنگار تجربی منحنی مدل پارامتری نمایی برازش داده شده (آ) برآوردگر کلاسیک، (ب) برآوردگر استوار
۲۶	۶.۲ نمودار سه بعدی نیم‌تغییرنگار و منحنی مدل پارامتری نمایی برازش داده شده در چهار جهت
۲۷	۷.۲ پیش‌گویی سه بعدی میدان نفتی تحت مطالعه
۲۸	۸.۲ نقشه پهنه بندی پیشگویی مقادیر مقاومت فشاری در عمق‌های ۱۷۰۳، ۲۱۵۵ و ۲۴۸۰ متر به ترتیب در (آ)، (ج) و (ه). همچنین نقشه پهنه‌بندی واریانس پیشگویی در همین عمق‌ها به ترتیب در (ب)، (د) و (و).
۲۹	۹.۲ مقادیر پیشگویی در مقابل مقادیر واقعی داده‌های آزمون
۳۹	۱.۳ شماتیک شبکه‌بندی به روش میانه‌پالایش (آ) و (ج) سه‌بعدی، (ب) دو بعد

۲.۳	نمودار پراکنش داده‌های خام و پراکنش نسبت به مختصات x ، y و z به ترتیب (آ)، (ب)، (ج) و (د).	۴۷
۳.۳	نمودار سه بعدی نیم‌تعییرنگار و (آ): منحنی مدل پارامتری خطی برازش داده شده روی داده‌های اصلی. (ب) منحنی مدل پارامتری نمایی برازش داده شده روی باقی‌مانده‌های حاصل از میانه‌پالایش. (ج) منحنی مدل پارامتری نمایی برازش داده شده روی باقی‌مانده‌های مدل برازش داده شده با استفاده از رگرسیون خطی	۴۸
۴.۳	مقادیر پیشگویی در مقابل مقادیر واقعی داده‌های آزمون با استفاده از کریگینگ میانه‌پالایش (آ) ۵ بسته، (ج) ۱۰ بسته و (ه) LOOCV و با استفاده از کریگینگ عام (ب) ۵ بسته، (د) ۱۰ بسته و (و) LOOCV.	۴۹
۵.۳	نمودار پراکنش داده‌ها و پراکنش نسبت به مختصات x ، y و z به ترتیب (آ)، (ب)، (ج) و (د).	۵۱
۶.۳	مقادیر پیشگویی در مقابل مقادیر واقعی داده‌های آزمون با استفاده از کریگینگ میانه‌پالایش (آ) ۵ بسته، (د) ۱۰ بسته و (ز) LOOCV، کریگینگ عام با روند خطی (ب) ۵ بسته، (ه) ۱۰ بسته و (ح) LOOCV و با استفاده از کریگینگ عام روی باقی‌مانده‌های حاصل از اسپلین (ج) ۵ بسته، (و) ۱۰ بسته و (ط) LOOCV.	۵۲
۷.۳	نمودار پراکنش داده‌ها نسبت به مختصات طول، عرض و عمق به ترتیب (آ)، (ب) و (ج).	۵۳
۸.۳	نمودار سه بعدی نیم‌تعییرنگار و منحنی مدل پارامتری نمایی برازش داده شده به باقی‌مانده‌های حاصل از روش‌های (آ): اسپلین، (ب): ماکسیمم درست‌نمایی و (ج): میانه‌پالایش.	۵۴
۹.۳	نمودار جعبه‌ای مقادیر MSPE حاصل از روش‌های ۱: کریگینگ میانه‌پالایش، ۲: کریگینگ عام و ۳: کریگینگ معمولی با استفاده از اسپلین.	۵۵
۱.۴	(آ) مکان مُد، و سیستم مختصات z ، (ب) نمایش نقاط منتخب انتگرال‌گیری با راهبرد توری	۶۴
۲.۴	نقاط منتخب انتگرال‌گیری با استفاده از (آ) روش CCD و (ب) روش توری	۶۷
۳.۴	موقعیت نقاط در طرح مرکب مرکزی	۶۸
۴.۴	توزیع نرمال استاندارد (خط ممتد) و چگالی معادله‌ی (۱۴.۴) برای مقادیر مختلف پارامتر مقیاس (خطوط خط‌چین)	۷۰
۵.۴	شماتیک ماتریسی معادله‌ی تاکاهاشی	۷۲
۶.۴	گسسته‌سازی ناحیه مثال شبیه‌سازی در فضای سه بعدی با چهاروجهی‌ها (آ) و یک چهاروجهی منتخب (ب)	۸۲
۷.۴	(آ) گسسته‌سازی سه بعدی ناحیه شبیه‌سازی شده، (ب) منحنی تابع کوواریانس واقعی در مثال شبیه‌سازی به همراه تقریب حاصل از روش SPDE	۸۳

فهرست جداول

۵۰		مقادیر MSPE	۱.۳
۵۰		مقادیر MSPE	۲.۳
۵۴		پارامترهای تغییرنگار	۳.۳
۵۵		پارامترهای تغییرنگار	۴.۳
۶۶		نقاط طرح هادامارد والش برای $k = 3$	۱.۴
۶۸		نقاط طرح هادامارد والش برای $k = 3$	۲.۴
۶۸		تعداد اجرا در طرح عاملی کسری	۳.۴
۸۴		برآوردهای میانه پسینی به همراه فواصل باورمند بیزی ۹۵ درصد	۴.۴

فصل ۱

آمار فضایی

ظهور علم آمار در داده‌هایی که به موقعیت مکانی و جهت قرارگیری از هم وابسته هستند، اولین بار در تحلیل نقشه‌ها رخ داد. برای مثال **هالی (۱۶۸۶)** جهت بادهای تجاری و موسمی مناطق استوایی را بر روی بخش‌هایی از نقشه‌ی زمین بررسی کرد. وی سعی کرد جهت این دو نوع باد را به یک عامل فیزیکی نسبت دهد. **استیودنت (۱۹۰۷)** به بررسی توزیع ذرات معلق درون مایعات علاقه‌مند بود. او به‌جای تحلیل مکانی ذرات، تعداد داده‌ها در هر ناحیه را به‌صورت گلبول‌شمار^۱ در هر میلی‌مترمربع شمارش و هر ناحیه را به چهارصد مربع تقسیم کرد و برای تحلیل از تعداد سلول‌های مخمر استفاده کرد. استیودنت متوجه شد که تعداد سلول‌های مخمر در هر مربع از توزیع پواسن تبعیت می‌کند. فیشر به وابستگی مکانی داده‌ها در آزمایش‌های میدانی کشاورزی واقف بود چرا که وی برای حذف این وابستگی‌ها، کرت‌های آزمایشی را در فواصل دور از هم برمی‌گزید. وی قوانین تصادفی‌سازی، بلوک‌بندی و تکراری را بنا کرد (**فیشر، ۱۹۳۵**). همچنین **یتز (۱۹۳۸)** با نشر یافته‌های پژوهش خود ادعا کرد که تصادفی‌سازی علاوه بر کنترل اریبی ناخواسته، اثر همبستگی مکانی را نیز کاهش می‌دهد، ولی به‌طور کامل از بین نمی‌برد. **اسمیت (۱۹۳۸)** به دنبال انتخاب ابعاد کرت‌های کشاورزی بود و دریافت که هرگونه افزایش در اندازه‌ی کرت، کاهش واریانس خطا را نتیجه می‌دهد. اگرچه تحلیل‌های او تجربی بودند و فرمول‌بندی آن‌ها دشوار اما وجود همبستگی مکانی را در آزمایشات میدانی‌اش به رسمیت می‌شناخت. همچنین از روش‌های نزدیک‌ترین

^۱Hemocytometer

همسایگی^۲ برای تحلیل داده‌ها در علوم کشاورزی با در نظر گرفتن همبستگی مکانی به‌طور غیرمستقیم، یا با استفاده از باقیمانده‌های همسایگی‌های کرت‌های مجاور به‌عنوان متغیرهای مورد بررسی استفاده می‌شد (پاپدکیس، ۱۹۳۷؛ بارتلت، ۱۹۷۸؛ ویلکینسون و همکاران، ۱۹۸۳؛ بساگ، ۱۹۸۶). در حال حاضر کاربرد آمار را تقریباً در هر شاخه علمی می‌توان یافت. در حیطه‌ی زمین‌شناسی، زیست‌شناسی و علوم زیست محیطی، غالباً (و البته نه همیشه) داده‌ها را نمی‌توان به‌صورت تصادفی در نظر گرفت یا بلوک‌بندی کرد و همچنین امکان تکرار داده‌ها وجود ندارد. پس نیاز به مدل‌های آماری بر اساس رویکردهای جدید برگرفته از تکنولوژی‌های قدیمی و جدید، وجود دارد. حل بسیاری از مسائل موجود از قبیل ارزیابی منابع، نظارت بر محیط زیست (به‌عنوان مثال، گرم شدن هوای کره زمین) و تصویربرداری پزشکی، تابع اطلاعات برگرفته از داده‌هایی است که معمولاً به موقعیت مکانی و جهت قرارگیری آن‌ها در ناحیه جغرافیایی مورد مطالعه بستگی دارد. برای تحلیل این نوع داده‌ها، که به داده‌های فضایی معروف هستند، روش‌های کلاسیک استنباط آماری جوابگو نیستند. وارد کردن ساختار وابستگی القاشده توسط موقعیت‌های مکانی داده‌ها (ساختار وابستگی فضایی) در تحلیل داده‌ها توسط روش‌های آمار فضایی ممکن شده است (کرسی، ۱۹۹۳).

در این بخش ابتدا انواع داده‌های فضایی را معرفی می‌کنیم و سپس مفاهیم اولیه آمار فضایی، پیشگوی فضایی و استنباط بیزی را ارائه می‌دهیم.

۱.۱ انواع داده‌های فضایی

در توصیف ماهیت داده‌های فضایی، تمایز بین گسستگی و پیوستگی فضای موقعیت‌ها و نیز مقادیر اندازه‌گیری شده برای هر متغیر مهم است. طبقه‌بندی داده‌های فضایی با نوع مفهوم فضا و سطح اندازه‌گیری اولین گام ضروری در انتخاب روش آماری مناسب برای حل یک سوال است. اما این طبقه‌بندی به تنهایی کافی نیست، چرا که ممکن است یک شی فضایی نمایان‌گر فضاهای جغرافیایی کاملاً متفاوتی باشد. (کرسی ۲۰۱۵) داده‌های فضایی را بر اساس تقسیم‌بندی سنتی و بسته به نوع مسئله به سه نوع داده‌های شبکه‌ای^۳، الگوی نقطه‌ای^۴ و زمین‌آمار^۵ تقسیم کرد. گاهی زمان به‌عنوان یک پارامتر در داده‌های فضایی تاثیرگذار است و با گذشت زمان داده‌ها اندازه‌گیری می‌شوند، از مدل‌های فضایی-زمانی^۶ می‌توان برای مدل‌بندی این نوع از داده‌ها، استفاده کرد (کرسی و ویکل، ۲۰۱۵). علاوه بر این فیشر و وانگ (۲۰۱۱) بر اساس یک نوع‌شناسی، داده‌های فضایی را به چهار گروه داده‌های

^۲Nearest-neighbor methods

^۳Lattice data

^۴Point pattern

^۵Geostatistical data

^۶spatio-temporal

زمین‌آماري، الگوی نقطه‌ای، ناحیه‌ای^۷ و داده‌های متقابل فضایی^۸ تقسیم کردند. متغیرهای اندازه‌گیری شده در آمار فضایی ممکن است گسسته یا پیوسته باشند. در ادامه هر کدام از انواع داده‌های فضایی را معرفی می‌کنیم.

داده‌های زمین‌آماري: این نوع از داده‌ها که داده‌های میدانی^۹ نیز نامیده می‌شوند، مربوط به متغیرهایی هستند که از لحاظ مفهومی پیوسته (دید میدانی) و مشاهدات آن‌ها از یک مجموعه ثابت و از پیش تعریف‌شده، نمونه‌برداری می‌شوند. در این نوع از داده‌ها متغیر مورد بررسی می‌تواند گسسته یا پیوسته باشد. برای حالت پیوسته می‌توان اندازه‌گیری مقدار مقاومت فشاری تک‌محوری سازند را در چاه مثال زد و برای حالت گسسته می‌توان تعداد کرم‌های خاکی درون نمونه خاک‌های برداشت‌شده در مکان‌های یک زمین کشاورزی حاصلخیز را نام برد.

داده‌های الگوی نقطه‌ای: به داده‌های متناهی در ناحیه‌ی مورد مطالعه، متشکل از مجموعه مکان‌های نقطه‌ای که رخداد مورد نظر در آن نقاط روی داده است، داده‌های الگوی نقطه‌ای می‌گویند. در این نوع از داده‌ها مکان خود یک متغیر است. داده‌های الگوی نقطه‌ای در سه دسته کاملاً فضایی تصادفی^{۱۰}، منظم^{۱۱} و خوشه‌ای^{۱۲} تقسیم‌بندی می‌شوند (کرسی، ۱۹۹۳). وجود یک نوع بیماری خاص در منطقه‌ی تحت مطالعه یا بروز یک نوع جرم را می‌توان مثالی برای این نوع داده‌ها دانست.

داده‌های ناحیه‌ای: داده‌های ناحیه‌ای یا شبکه‌ای^{۱۳} داده‌هایی متناظر با مشاهدات با تعداد ثابت از یک فضای واحد (شی فضایی) هستند که ممکن است یک شبکه منظم مانند تصاویر سنجش از دور یا مجموعه‌ی نامنظم منطقه یا محدوده، مانند بخش و شهرستان‌های محدوده‌ی سرشماری یا حتی کشورها باشد. هدف اصلی از تحلیل داده‌های فضایی ناحیه‌ای مدل‌بندی احتمالاتی مشاهدات است، در حالی که در زمین‌آمار هدف اصلی پیش‌گویی مقدار متغیر در یک مکان بدون مشاهده است.

داده‌های متقابل فضایی: این نوع از داده‌ها به داده‌های روند مبدأ-مقصد یا داده‌های ربطی شهرت دارند و مقدار اندازه‌گیری شده مربوط به پیوند بین دو نقطه یا دو ناحیه را شامل می‌شود. داده‌های مربوط به تجارت بین‌المللی یا مهاجرت از جمله داده‌های متقابل فضایی هستند.

^۷ Areal

^۸ Spatial interaction data

^۹ Field data

^{۱۰} Complete spatial randomness

^{۱۱} Regularity

^{۱۲} Clustering

^{۱۳} Lattice data

۲.۱ مدل بندی آماری فضایی

برای تحلیل داده‌های فضایی، یک مدل آماری در نظر گرفته می‌شود. در آمار فضایی معمولاً یک میدان تصادفی که مجموعه‌ای از متغیرهای تصادفی مانند $\{Z(s); s \in D\}$ است به عنوان مدل آماری داده‌های فضایی در نظر گرفته می‌شود، که در آن مجموعه اندیس گذار D یک زیر مجموعه از فضای اقلیدسی d بعدی، $d \geq 1$ ، از $\mathbb{R}^d$ است. در مورد میدان تصادفی $Z(\cdot)$ ، میانگین در موقعیت s و کوواریانس در موقعیت‌های s_1 و s_2 به ترتیب به صورت

$$\begin{aligned}\mu(s) &= E[Z(s)], s \in D \\ C(s_1, s_2) &= Cov[Z(s_1), Z(s_2)] \\ &= E[(Z(s_1) - \mu(s_1))(Z(s_2) - \mu(s_2))], s_1, s_2 \in D\end{aligned}\quad (1.1)$$

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

$$Var[Z(s)] = E[(Z(s) - \mu(s))^2] = C(s, s) \quad (2.1)$$

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

$$Z(s) = \mu(s) + \delta(s), \quad s \in D \quad (3.1)$$

تجزیه کرد، که در آن $\mu(\cdot)$ تغییرات بزرگ مقیاس^{۱۴} یا روند^{۱۵} و $\delta(\cdot)$ فرآیند خطا یا تغییرات کوچک مقیاس^{۱۶} میدان نامیده می‌شوند. تغییرات کوچک مقیاس ممکن است ناشی از خطای اندازه‌گیری یا تغییرپذیری در درون موقعیت مشاهده شده باشد و تغییرات بزرگ مقیاس ممکن است ناشی از تغییرپذیری بین موقعیت‌های مشاهده شده باشند. به طور معمول تحلیل داده‌های فضایی بر اساس مشاهدات نمونه بسیار دشوار است، اما برخی ویژگی‌های میدان تصادفی موجب ساده‌سازی مسأله خواهند شد، که در ادامه معرفی می‌شوند (محمدزاده، ۱۳۹۸).

تعریف ۱.۲.۱. میدان تصادفی $Z(\cdot)$ ، مانای ذاتی^{۱۷} نامیده می‌شود، اگر

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

$$2. \text{ برای عبارت } (Z(s_i) - Z(s_j)) \text{ واریانس باید تنها تابعی از موقعیت‌های } s_i \text{ و } s_j \text{ باشد، یعنی}$$

$$Var(Z(s_i) - Z(s_j)) = 2\gamma(s_i - s_j), \quad s_i, s_j \in D \subset \mathbb{R}^d.$$

بنابراین واریانس به مکان قرارگیری مشاهدات از هم بستگی ندارد و در همه جای میدان ثابت است.

^{۱۴} Large scale variation

^{۱۵} Trend

^{۱۶} Small scale variation

^{۱۷} Intrinsic stationary

تعریف ۲.۲.۱. میدان تصادفی $Z(\cdot)$ ، مانای مرتبه دوم^{۱۸} یا مانای ضعیف نامیده می‌شود، اگر

۱. میانگین میدان تصادفی $Z(\cdot)$ ، مستقل از مکان s یا ثابت باشد. یعنی

$$E(Z(s)) = \mu, \quad s \in D \subset \mathbb{R}^d.$$

۲. کوواریانس $Z(s_i)$ و $Z(s_j)$ فقط تابعی از موقعیت‌های s_i و s_j باشند، یعنی

$$\text{Cov}(Z(s_i) - Z(s_j)) = C(s_i - s_j), \quad s_i, s_j \in D \subset \mathbb{R}^d.$$

واضح است که تحت مانایی مرتبه دوم

$$\text{Var}(Z(s)) = \text{Cov}(Z(s) - Z(s)) = C(o) = \sigma^2.$$

تعریف ۳.۲.۱. میدان تصادفی $Z(\cdot)$ ، مانای قوی^{۱۹} نامیده می‌شود، هرگاه برای همه‌ی موقعیت‌های

$s_1, \dots, s_n$ و گام h توزیع توأم $Z(s_1), \dots, Z(s_n)$ با توزیع توأم $Z(s_1 + h), \dots, Z(s_n + h)$ یکی باشد، یعنی

$$(Z(s_1), \dots, Z(s_n)) \stackrel{D}{=} (Z(s_1 + h), \dots, Z(s_n + h)).$$

به عبارتی، در صورت انتقال موقعیت‌های $s_1, \dots, s_n$ در راستای $h \in \mathbb{R}^d$ توزیع توأم $Z(s_1), \dots, Z(s_n)$ تغییر نمی‌کند.

می‌توان ثابت کرد که مانای قوی شرط کافی برای مانای مرتبه دوم و مانای مرتبه دوم شرط کافی برای مانای ذاتی است، ولی در حالت کلی عکس آن برقرار نیست (محمدزاده، ۱۳۹۸).

۳.۱ ساختار همبستگی فضایی

در آمار فضایی از تغییرنگار^{۲۰} و هم‌تغییرنگار^{۲۱} برای مدل‌بندی ساختار همبستگی داده‌های فضایی استفاده می‌شود. در اکثر موارد داده‌های نزدیک به هم شباهت بیشتری دارند. پس متوسط تفاضل مقادیر $Z(s)$ و $Z(s+h)$ که به فاصله‌ی h از یکدیگر قرار دارند، می‌تواند معیاری مناسب برای بیان اندازه‌ی این تشابه و وابستگی متقابل مقادیر میدان تصادفی در دو موقعیت s و $s+h$ باشد. هر چه متوسط این تفاضل کوچک باشد، واریانس میدان تصادفی کوچک است و می‌توان گفت $Z(s)$ و $Z(s+h)$ به یکدیگر نزدیکتر و به عبارتی شباهت بیشتری به هم دارند. به همین صورت هر چه متوسط تفاضل‌ها بیشتر باشد، واریانس مقادیر میدان تصادفی بزرگتر بوده و تشابه آن‌ها کمتر است.

^{۱۸}Second order stationary

^{۱۹}Strong stationary

^{۲۰}Variogram

^{۲۱}Covariogram

۱.۳.۱ تغییرنگار

تغییرنگار نظری، $\gamma(h)$ ، واریانس تفاضل بین دو مقدار مشاهده‌شده میدان تصادفی در دو مکان فضایی s و $s+h$ است که به فاصله h از هم قرار دارند و به‌صورت زیر تعریف می‌شود:

$$\gamma(h) = Var(Z(s+h) - Z(s)). \quad (۴.۱)$$

تابع $\gamma(h)$ نیم‌تغییرنگار^{۲۲} نامیده می‌شود. این تابع برای توصیف میزان وابستگی مقادیر میدان تصادفی و برای پیش‌گویی قابل اعتماد توسط پیشگوی فضایی معروف به کریگینگ^{۲۳} مورد نیاز است. تغییرنگار تجربی (نمونه) معمولاً توسط روش‌های گشتاوری در دنباله‌ای از گام‌ها محاسبه می‌شود و یک یا چند مدل معتبر (مجاز^{۲۴}) تغییرنگار بر روی آن برازش داده می‌شود. یک تغییرنگار ممکن است در طول یک برش یا توری در فواصل منظم یا در فواصل پراکنده نامنظم محاسبه شود. برآوردهای تغییرنگار ممکن است به‌صورت گشتاوری محاسبه شوند یا این که برای محاسباتشان از میانه و مد استفاده شود که برآوردهای تنومند^{۲۵} نامیده می‌شوند و نسبت به مقادیر دورافتاده^{۲۶} حساس نیستند. تغییرنگار ممکن است کراندار (در یک میدان مانای مرتبه دوم) یا بدون کران (در یک میدان مانای ذاتی) باشد. پارامترهای تغییرنگار که توسط مدل‌های معتبر تغییرنگار برآورد می‌شوند، خلاصه‌ای از تغییرات فضایی را که برای پیشگویی فضایی نیاز است به ما می‌دهند. میانگین کمترین توان‌های دوم باقی‌مانده‌ها^{۲۷} و معیار آکاییک^{۲۸} می‌تواند در انتخاب بهترین مدل برازش داده شده به ما کمک کند (اولیور و وستر، ۲۰۱۵). برآورد تغییرنگار معمولاً براساس روش گشتاوری ماترون^{۲۹} (MOM) به‌صورت

$$\hat{\gamma}(h) = \frac{1}{2m(h)} \sum_{i=1}^{m(h)} \{z(s_i) - z(s_i + h)\}^2 \quad (۵.۱)$$

محاسبه می‌شود، که در آن $z(s_i)$ و $z(s_i + h)$ مقادیر مشاهده‌شده میدان $Z(\cdot)$ در موقعیت‌های s_i و $s_i + h$ هستند و $m(h)$ تعداد جفت نمونه‌هایی است که در فاصله h از هم قرار دارند. اگر تغییرنگار برای تمام مقادیر h تقریباً ثابت باشد، داده‌ها همبستگی فضایی ندارند یا بسیار ضعیف است. برعکس، شیب نمودار تغییرنگار بیانگر همبستگی فضایی داده‌ها است. از عوامل موثر بر تغییرنگار تجربی می‌توان اندازه نمونه، فاصله‌ی نمونه‌گیری و مقیاس فضایی، فاصله‌ی گام‌ها از هم، توزیع داده‌ها، همسان‌گردی^{۳۰} و روند نام برد (اولیور و وستر،

^{۲۲} Semivariogram

^{۲۳} Kriging

^{۲۴} Authorized

^{۲۵} Robust estimators

^{۲۶} Outliers

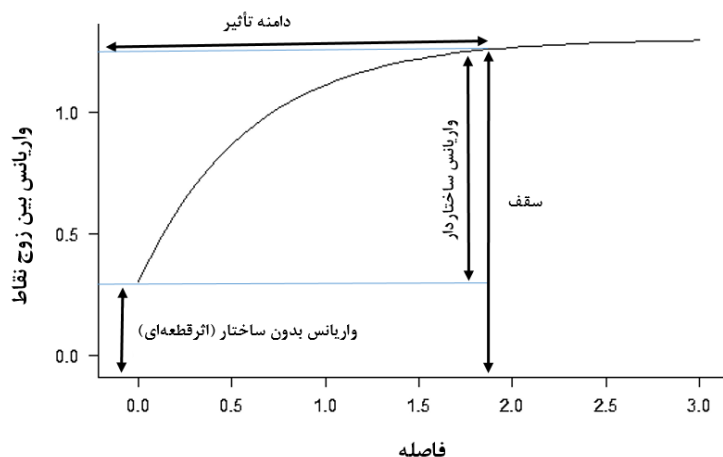
^{۲۷} Mean squared residuals

^{۲۸} Akaike

^{۲۹} Matheron's method of moments

^{۳۰} Isotropy

۲۰۱۵). تغییرنگار دارای سه پارامتر اثرقطعه‌ای^{۳۱} (واریانس بدون ساختار)، سقف^{۳۲} (آستانه یا ازاره) و دامنه^{۳۳} است که به صورت زیر تعریف می‌شوند.



شکل ۱.۱: اجزاء تشکیل دهنده تغییرنگار

- **دامنه:** در آمار فضایی قاعدتاً مشاهدات نزدیک به هم تشابه بیشتری نسبت به نقاط دور از هم دارند، لذا با افزایش فاصله بین دو مشاهده، مقدار تغییرنگار افزایش می‌یابد. افزایش تغییرنگار تا یک فاصله‌ی معین ادامه دارد و سپس به ازاء مقادیر بزرگتر فاصله مقدار تغییرنگار ثابت می‌ماند. به عبارتی نمونه‌ها تا فاصله‌ی معینی بر هم تأثیر می‌گذارند و دارای وابستگی مکانی هستند. فاصله‌ای که در آن تغییرنگار به حد ثابتی می‌رسد و به خط افقی نزدیک می‌شود را دامنه یا شعاع تأثیر می‌نامند. دامنه تأثیر بزرگتر بر وابستگی مکانی گسترده‌تر دلالت دارد. خارج از دامنه‌ی تأثیر مشاهدات بر هم اثری ندارند و با استفاده از آمار کلاسیک می‌توان این داده‌ها را تحلیل کرد.

- **سقف:** به مقدار ثابتی که تغییرنگار در دامنه‌ی تأثیر به آن می‌رسد، سقف یا آستانه گفته می‌شود. مقدار سقف برابر با واریانس تمام نمونه‌هایی است که در محاسبه‌ی تغییرنگار به کار رفته است. تغییرنگارهایی که به سقف مشخص می‌رسند به اصطلاح کراندار هستند و از اهمیت بیشتری در پیشگویی برخوردارند. در مواردی ممکن است که تغییرنگار کراندار نباشد که می‌تواند نشان دهنده‌ی وجود روند یا عدم نایستایی داده‌ها باشد.

- **اثرقطعه‌ای:** به لحاظ نظری مقدار تغییرنگار در مبداء مختصات، $h = 0$ ، باید برابر با صفر باشد. ولی در عمل، تغییرنگارهای واقعی که از تجربه حاصل می‌شوند از چنین دستوری تبعیت نمی‌کنند. به مقدار تغییرنگار به ازاء $h = 0$ اثرقطعه‌ای می‌گویند که با C_0 نمایش

^{۳۱}Nugget effect

^{۳۲}Sill

^{۳۳}Range

می‌دهند. **کرسی** (۱۹۹۳) دو خطای نمونه‌گیری و خطای اندازه‌گیری یا تغییرات شدید کمیت را دلایل صفر نبودن اثر قطعه‌ای می‌شمارد و به صورت کلی اثر قطعه‌ای را جمع دو خطای اندازه‌گیری^{۳۴} (c_{ME}) و خطای ریزمقیاس^{۳۵} (c_{MS}) می‌داند. در منابع زمین‌آماري عوامل تاثیرگذار بر مقدار بیشتر از صفر اثر قطعه‌ای را به دو دسته تقسیم کرده‌اند:

۱. وجود مولفه‌های تصادفی در توزیع متغیر، که در واقع به تصادفی بودن متغیرها برمی‌گردد.

۲. خطاهای نمونه‌برداری، آماده‌سازی و تحلیل آزمایشگاهی (**حسني‌پاک**، ۱۳۷۷).

تفاضل مقدار سقف و اثر قطعه‌ای را واریانس ساختاردار، C ، می‌نامند که تابعی از موقعیت مکانی و جهت قرارگیری داده‌ها است. واریانس ساختاردار تغییراتی است که می‌توان دلیل آن را در خصوصیات خود متغیر فضایی یافت. از نسبت $\frac{C}{C+C_0}$ می‌توان برای تشخیص وجود ساختار فضایی در داده‌ها استفاده کرد، که برابر است با نسبت مولفه‌ی ساختاردار به مولفه‌ی بدون ساختار تغییرنگار. نسبت دیگری نیز بدین منظور وجود دارد که برابر است با $\frac{C_0}{C+C_0}$ و معرف آن است که چه مقدار از تغییرپذیری را اثر قطعه‌ای توجیه می‌کند. در صورتی که $\frac{C}{C+C_0} > \frac{1}{4}$ باشد، نقش مولفه غیرساختارمند از مولفه ساختاردار کمتر است و می‌توان نتیجه گرفت ساختار فضایی داده‌ها قابل ملاحظه است. هر چه مقدار قدرنسبت فوق به یک نزدیکتر باشد استفاده از آمار فضایی ضرورت بیشتری پیدا می‌کند.

۲.۳.۱ هم‌تغییرنگار

کوواریانس دو مقدار $Z(s)$ و $Z(s+h)$ به صورت

$$C(h) = \text{Cov}(Z(s), Z(s+h)) \quad (۶.۱)$$

تعریف می‌شود و آن را هم‌تغییرنگار^{۳۶} می‌نامند.

اگر میدان تصادفی، مانای مرتبه‌ی دوم باشد آن‌گاه رابطه‌ی (۶.۱) به صورت

$$C(h) = E[Z(s)Z(s+h)] - \mu^2$$

ساده می‌شود. تابع هم‌تغییرنگار، تابعی همیشه مثبت^{۳۷} است (**کرسی**، ۱۹۹۳). یعنی برای هر دنباله از اعداد حقیقی ثابت $\{a_i\}_{i=1}^m$ ، $\sum_{i=1}^m \sum_{j=1}^m a_i a_j C(s_i - s_j) \geq 0$. هم‌تغییرنگار نمایشگر میزان شباهت تغییرپذیری دو متغیر $Z(s)$ و $Z(s+h)$ است.

^{۳۴} Measurment error

^{۳۵} Micro scale error

^{۳۶} Covariogram

^{۳۷} Positive definite

۳.۳.۱ همبستگی نگار

ضریب همبستگی بین $Z(s)$ و $Z(s+h)$ همبستگی نگار^{۳۸} نام دارد و با شرایط $C(0) > 0$ به صورت

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

تعریف می شود.

با توجه به این که تغییرنگار تابعی از فاصله است، همبستگی نگار نیز تابعی از فاصله خواهد بود. بنابراین با افزایش فاصله بین موقعیت ها، مقدار همبستگی نگار کم و با کاهش فاصله مقدار آن افزایش پیدا می کند.

۴.۳.۱ ناهمسان گردی

تابع تغییرنگار ممکن است به موقعیت مکانی و به طور نسبی به جهت جغرافیایی قرارگیری داده ها در ناحیه مورد مطالعه وابسته باشد. اگر این تابع فقط به فاصله مکانی داده ها از یکدیگر بستگی داشته باشد و به خود مکان و جهت جغرافیایی وابسته نباشد، میدان تصادفی فضایی حاصل را، که برای مدل بندی داده ها استفاده می شود، همسان گرد^{۳۹} و در غیر این صورت ناهمسان گرد^{۴۰} می نامند (زیمرمن، ۱۹۹۳؛ مونتر و همکاران، ۲۰۱۵). پذیره همسان گردی میدان تصادفی موجب آسان تر شدن مدل بندی داده ها، برآورد تغییرنگار و هم تغییرنگار و در نهایت افزایش دقت در نتایج حاصل از تحلیل داده های فضایی مثل پیش گویی می شود. بنابراین، بررسی همسان گردی ساختار همبستگی داده ها قبل از تحلیل آن ها یک ضرورت محسوب می شود (محمدزاده، ۱۳۹۸). حسنی پاک (۱۳۷۷) و سارما (۲۰۱۰) دو نوع ناهمسان گردی هندسی^{۴۱} و ناهمسان گردی ناحیه ای^{۴۲} را مورد بحث قرار داده اند. همچنین زیمرمن (۱۹۹۳) به بررسی ناهمسان گردی دامنه ای^{۴۳} و ناهمسان گردی اثر قطعه ای^{۴۴} پرداخت. علاوه بر این ها، اریکسون و سیسکا (۲۰۰۰) ناهمسان گردی در شیب و توان را در توابع تغییرنگار مورد بررسی قرار داده اند. در این جا سعی کرده ایم سه نوع ناهمسان گردی دامنه ای، ازاره ای و اثر قطعه ای را به اختصار تشریح کنیم.

● **ناهمسان گردی هندسی:** مهم ترین نوع ناهمسان گردی است که به آن ناهمسان گردی دامنه ای یا ناهمسان گردی بیضوی نیز گفته می شود (مونتر و همکاران، ۲۰۱۵). اگر

^{۳۸}Correlogram

^{۳۹}Isotropic

^{۴۰}Anisotropic

^{۴۱}Geometric Anisotropy

^{۴۲}Zonal Anisotropy

^{۴۳}Range Anisotropy

^{۴۴}Nugget Anisotropy

تغییرنگار میدان در جهت‌های مختلف دارای ازاره یکسان ولی دامنه تاثیر متفاوتی باشد، ناهمسان‌گردی هندسی رخ داده است.

- **ناهمسان‌گردی اثر قطعه‌ای:** اگر مقدار اثر قطعه‌ای تغییرنگار در جهت‌های مختلف یکسان نباشد، این نوع ناهمسان‌گردی را ناهمسان‌گردی اثر قطعه‌ای می‌نامند.
- **ناهمسان‌گردی ازاره‌ای:** اگر مقدار ازاره (سقف) تغییرنگار در جهت‌های مختلف یکسان نباشد، ناهمسان‌گردی ازاره‌ای رخ داده است. در منابع زمین‌آماري این نوع ناهمسان‌گردی را ناهمسان‌گردی ناحیه‌ای نیز می‌نامند (حسنی‌پاک، ۱۳۷۷).

تعریف ۱.۳.۱. میدان تصادفی‌ای که مانای مرتبه دوم و همسان‌گرد باشد، همگن^{۴۵} نامیده می‌شود.

۵.۳.۱ خانواده‌های توابع همبستگی فضایی

این بحث را با این سوال مطرح می‌کنیم که چه ویژگی‌هایی را می‌توان برای خانواده‌های توابع همبستگی فضایی انتظار داشت؟ در ابتدا اگر ما به قانون جغرافیایی توبلر^{۴۶} باور داشته باشیم، انتظار داریم که با افزایش فاصله مکانی s_i و s_j همبستگی فضایی بین $Z(s_i)$ و $Z(s_j)$ کاهش یابد. به عبارتی $\rho(h)$ یک تابع غیرافزایشی از h است. همچنین، با استفاده از واحدهای مختلف می‌توان فاصله را اندازه‌گیری کرد، پس $\rho(h)$ باید دارای یک پارامتر مقیاس‌گذاری واحد باشد (دیگل و جورجی، ۲۰۱۹). با این مقدمه برخی از خانواده‌های پارامتری که این دو ویژگی را دارند، عبارتند از خانواده‌های نمایی، مترن و کروی که در بخش بعدی مورد بررسی قرار گرفته‌اند.

خانواده نمایی

تابع همبستگی معتبر نمایی به صورت (۷.۱) تعریف می‌شود:

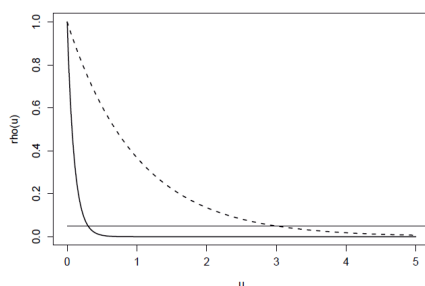
$$\rho(h, \phi) = \exp(-h/\phi) : h \geq 0, \quad (7.1)$$

که در آن $\phi > 0$ تنها پارامتر مدل است. شکل ۲.۱ مقدار $\rho(h, \phi)$ را به ازاء $\phi = 1^\circ$ و $\phi = 1/2^\circ$ نشان می‌دهد. پارامتر مقیاس‌بندی ϕ را پارامتر دامنه می‌نامند. برای واحد مشخصی از فاصله، مقدار بزرگتری از ϕ نشان می‌دهد که مقدار همبستگی فضایی از دامنه بیشتر است. یک تعریف مرسوم از دامنه عملی^{۴۷} در همبستگی فضایی فاصله‌ای است که به ازای آن $\rho(h) = 0.05$ شود. برای تابع همبستگی نمایی (۷.۱) دامنه عملی تقریباً سه برابر ϕ است.

^{۴۵}Homogeneous

^{۴۶}Tobler

^{۴۷}Practical range



شکل ۲.۱: تابع همبستگی نمایی به ازاء $\phi = 0.1$ (خط ممتد) و $\phi = 1$ (خط چین). محل تقاطع خط افقی باریک ($\rho(h) = 0.05$) با منحنی مقدار دامنه‌ی عملی را نشان می‌دهد.

خانواده مترن

پارامتر دامنه ϕ در مدل نمایی کنترل مرتبه‌ی سرعت میل کردن همبستگی فضایی به سمت صفر را با استفاده از جداسازی ممکن می‌سازد. خانواده مترن این انعطاف‌پذیری را به شکل تابع همبستگی می‌افزاید و به نوبه‌ی خود تاثیر مستقیمی بر همواری فرآیند فضایی دارد. مترن (۱۹۱۷-۲۰۰۷) آماردان سوئدی شاغل در دانشگاه سلطنتی جنگل‌داری استکهلم بود. او در رساله‌ی دکتری خود در سال ۱۹۶۰ نگاهی ژرف به روش‌های توسعه‌یافته آمار فضایی دارد که ماحصل آن را می‌توان معرفی خانواده دو پارامتری از توابع همبستگی دانست که با نام وی به ثبت رسیده است.

تابع همبستگی مترن به صورت

$$\rho(\mathbf{h}; \phi, k) = \{\Psi^{k-1} \Gamma(k)\}^{-1} (\mathbf{h}/\phi)^k K_k(\mathbf{h}/\phi)$$

است که در آن $\phi > 0$ و $k > 0$ پارامترهای مدل هستند، $\Gamma(\cdot)$ تابع گاما و $K_k(\cdot)$ تابع بسل تعدیل شده^{۴۸} نوع سوم از مرتبه‌ی k است. شکل ۳.۱ برای $k = (0.5, 1.5, 2.5)$ مقدار $\rho(\mathbf{h}; \phi, k)$ را نشان می‌دهد. پارامتر ϕ طوری انتخاب شده است که دامنه‌ی عملی برابر با سه باشد. تابع کوواریانس مترن به صورت

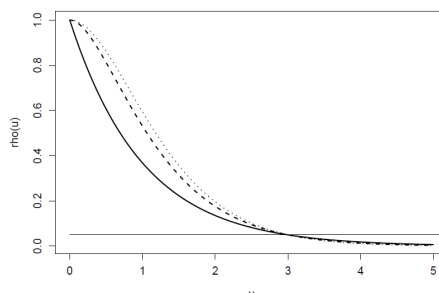
$$\text{Cov}(\mathbf{h}) = \frac{\sigma^2}{\Psi^{k-1} \Gamma(k)} (\phi \|\mathbf{h}\|)^k K_k(\phi \|\mathbf{h}\|) \quad (A.1)$$

که در آن $\|\mathbf{h}\|$ فاصله‌ی اقلیدسی بین s_i و s_j است، به طوری که $\|\mathbf{h}\| = s_i - s_j$ ، $K_k(\cdot)$ تابع بسل اصلاح‌شده‌ی نوع دوم از مرتبه‌ی k است و k درجه همواری تابع را کنترل می‌کند که اغلب مقدار ثابتی است.

خانواده کروی

تجربه نشان داده است که خانواده مترن به دلیل انعطاف‌پذیری و دامنه‌ی وسیعی که دارد می‌تواند برای بسیاری از مدل‌سازی‌های همبستگی فضایی کافی باشد. ولی در زمین‌شناسی

^{۴۸}Modified Bessel function



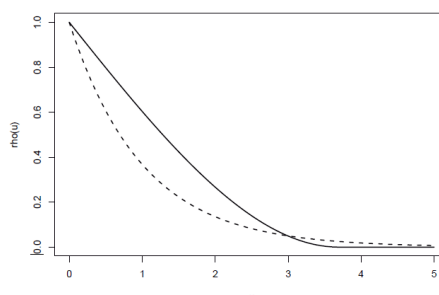
شکل ۳.۱: خانواده همبستگی مترن. برای $k = (0/5, 1/5, 2/5)$ (خط ممتد، خط چین و نقطه‌چین به ترتیب) مقدار ϕ طوری انتخاب شده که در هر سه مورد مقدار دامنه‌ی عملی برابر با ۳ شود و این مقدار از تقاطع خط نازک افقی ($\rho = 0/5$) با منحنی به دست می‌آید.

کلاسیک از مدل‌های دیگری نیز استفاده می‌شود که در اینجا فقط به شرح خانواده کروی که کاربرد بیشتری دارد، می‌پردازیم.

تابع همبستگی کروی از خانواده تک پارامتری است و به شکل

$$\rho(h) = \begin{cases} 1 - \frac{3}{4}(h/\phi) + \frac{1}{4}(h/\phi)^3 & : h \leq \phi \\ 0 & : h > \phi \end{cases} \quad (9.1)$$

تعریف می‌شود. در شکل ۴.۱ منحنی دو تابع همبستگی کروی و مترن برای پارامتر همواری $k = 1/5$ رسم شده‌اند.



شکل ۴.۱: تابع همبستگی کروی (خط ممتد) و مترن (خط چین) برای مقدار $k = 1/5$. در هر دو مورد مقدار ϕ طوری انتخاب شده است که دامنه‌ی عملی برابر با ۳ شود.

۴.۱ استنباط آماری

استنباط آماری فرآیند نتیجه‌گیری از داده‌ها است. به‌طور کلی مباحث استنباط شامل آزمایش کردن^{۴۹}، برآورد^{۵۰} و پیش‌گویی^{۵۱} است. استنباط و مدل‌بندی آماری رشد چشم‌گیری به‌ویژه برای مقابله با تغییرات و عدم قطعیت‌ها داشته‌اند. این موضوع یک امر بلندپروازانه به نظر می‌رسد، زیرا هر کسی که زندگی را پشت‌سر بگذارد حتی با کمترین تنش‌ها، عدم قطعیت‌های همه‌جانبه زندگی را به‌طور کامل درک می‌کند. بدیهی نیست که در مواجهه شدن با عدم قطعیت چیزی سخت‌گیرانه، علمی یا حتی منطقی بگوییم.

مکاتب مختلف فکری رشته آمار در واکنش به عدم قطعیت‌ها پدید آمده‌اند. دو مکتب غالب در دنیای امروز، استنباط مبتنی بر درست‌نمایی و استنباط بیزی هستند. در دیدگاه بیزی، تمامی عدم قطعیت‌ها می‌توانند از طریق قوانین استاندارد شده احتمال^{۵۲} مدل‌سازی و پردازش شوند (پاوتن، ۲۰۰۱). در ادامه به‌طور خلاصه این دو دیدگاه استنباطی را معرفی می‌کنیم.

۱.۴.۱ استنباط مبتنی بر درست‌نمایی

تعریف ۱.۴.۱. فرض کنید $f(z|\theta)$ تابع چگالی (جرم) احتمال نمونه $z = (z_1, \dots, z_n)$ باشد، به‌طوری که $Z = z$ مشاهده شده باشد. تابعی از θ که به‌صورت $L(\theta|z) = f(z|\theta)$ باشد را تابع درست‌نمایی^{۵۳} می‌نامند. برای مشاهدات مستقل، تابع درست‌نمایی برابر است با

$$L(\theta; z) = \prod_{i=1}^n f(z_i|\theta_1, \dots, \theta_k) \quad (10.1)$$

که در آن $\theta = (\theta_1, \dots, \theta_k)^T$ بردار پارامترهای نامعلوم است. در تابع چگالی کمیت z متغیر و مقدار θ ثابت است، اما در تابع درست‌نمایی مقدار θ می‌تواند تمامی مقادیر روی پارامتر را اختیار کند و z یک مقدار مشاهده شده است.

قضیه ۱.۴.۱. فرض کنید z و y دو مشاهده از فضای نمونه باشند، به‌طوری که تابع درست‌نمایی z متناسب با تابع درست‌نمایی y باشد، یعنی ثابتی مانند $C(z, y)$ داشته باشیم به‌طوری که برای تمام θ ‌ها داشته باشیم

$$L(\theta|z) = C(z, y)L(\theta|y).$$

^{۴۹}Testing

^{۵۰}Estimation

^{۵۱}Prediction

^{۵۲}Standard rules of probability

^{۵۳}Likelihood function

در این صورت نتایج استنباط برای پارامتر θ با استفاده از z مشابه نتایج حاصل با استفاده از y است. به این قضیه اصل درستنمایی^{۵۴} می‌گویند که در آن مقدار ثابت $C(z, y)$ به پارامتر θ بستگی ندارد.

نتیجه عملی استفاده از اصل درستنمایی، پایه‌ریزی استنباط‌های آماری بر اساس تابع درستنمایی است. یکی از محصولات این نتیجه، محاسبه برآوردهای ماکسیمم درستنمایی برای پارامترهای نامعلوم است.

تعریف ۲.۴.۱. فرض کنید $L(\theta|z)$ تابع درستنمایی برای متغیرهای تصادفی با مشاهدات $z_1, \dots, z_n$ باشد. اگر $\hat{\theta}$ مقدار تابع درستنمایی را بیشینه کند، آن‌گاه $\hat{\theta}$ برآوردهای درستنمایی ماکسیمم^{۵۵} θ تعریف می‌شود (**مود و همکاران، ۲۰۱۷**). برآورد درستنمایی ماکسیمم مقداری از فضای نمونه θ را اختیار می‌کند که در آن نقطه، z بیشترین شانس مشاهده شدن را داشته باشد. اگر تابع درستنمایی بر حسب پارامتر مشتق‌پذیر باشد، برآوردهای درستنمایی ماکسیمم در صورت وجود از حل معادله‌ی زیر به‌دست می‌آید:

$$\frac{\partial L(\theta|z)}{\partial \theta_i} = 0, \quad i = 1, \dots, k.$$

همچنین $L(\theta|z)$ و $l(\theta|z) = \ln L(\theta|z)$ در یک نقطه از θ ماکسیمم می‌شوند و در بعضی موارد، پیدا کردن ماکسیمم لگاریتم درستنمایی ساده‌تر است.

در مدل‌های واقعی ولی پیچیده آماری معمولاً محاسبه تابع درستنمایی کار مشکل و از نظر محاسباتی دارای پیچیدگی‌های مضاعفی است. معمولاً این مسأله، آماردانان را به استفاده از دیدگاه‌های استنباطی دیگر سوق می‌دهد.

۲.۴.۱ استنباط بیزی

رویکرد بیزی^{۵۶} اساساً با رویکرد کلاسیک و از آن جمله مبتنی بر درستنمایی متفاوت است. برخی جنبه‌های آمار بیزی می‌تواند بر سایر رویکردهای آماری مفید واقع شود. در آمار کلاسیک پارامتر نامعلوم و یک کمیت ثابت در نظر گرفته می‌شود، به‌طوری‌که نمونه‌های تصادفی $Z_1, \dots, Z_n$ از جامعه با شاخص θ انتخاب شده و براساس مقادیر مشاهده شده در نمونه، اطلاعاتی درباره‌ی پارامتر به ما می‌دهند. در رویکرد بیزی θ یک کمیت است که می‌تواند دارای توزیع احتمال (که توزیع پیشین^{۵۷} نامیده می‌شود) باشد. سپس اطلاعات نمونه برگرفته از جامعه‌ای با شاخص θ و توزیع پیشین باهم ترکیب شده و به اصطلاح به‌روز می‌شوند. این

^{۵۴} Likelihood principle

^{۵۵} Maximum likelihood estimation

^{۵۶} Bayesian approach

^{۵۷} Prior distribution

پیشین به روز شده را توزیع پسین^{۵۸} می نامند که با استفاده از قاعده ی بیز^{۵۹} انجام می شود. از این رو نام آمار بیزی را بر آن نهاده اند (کسلا و برگر، ۲۰۰۲). هر چه در توزیع پسین از اطلاعات پیشین پارامتر بیشتر استفاده شود، تاثیر پیشین در توزیع پسین بیشتر شده و به اصطلاح توزیع پیشین آگاهی بخش تر^{۶۰} خواهد بود. با این وجود اگر اطلاعات توزیع پیشین در دسترس نبوده یا مبهم باشد، از توزیع پیشین ناآگاهی بخش^{۶۱} یا مبهم^{۶۲} برای تحلیل بیزی استفاده می شود. توزیع های پیشین ناآگاهی بخش اغلب ناسره^{۶۳} هستند، یعنی

$$\int \pi(\theta) d\theta = \infty.$$

از طرفی چون توزیع پیشین ناسره توزیع احتمال نیستند، لذا ممکن است توزیع پسین سره را نتیجه ندهد و باید توجه داشت که تنها در صورت سره بودن توزیع پسین استنباط بیزی ممکن خواهد شد.

تعریف ۳.۴.۱. فرض کنید $Z_1, \dots, Z_n$ یک نمونه ی تصادفی با توزیع معلوم $\pi(z|\theta)$ باشد، به طوری که $\theta \in \Theta$ پارامتر مجهول باشد. اگر توزیع پیشین $\pi(\theta)$ مشخص باشد، آن گاه توزیع پسین، توزیع شرطی θ به شرط نمونه مشاهده شده $Z = z$ است و به صورت

$$\pi(\theta|z) = \frac{\pi(z|\theta)\pi(\theta)}{m(z)} \quad (11.1)$$

تعریف می شود، که در آن $m(z)$ توزیع کناری z و برابر است با

$$m(z) = \int \pi(z|\theta)\pi(\theta) d\theta. \quad (12.1)$$

نکته ۱.۴.۱. توزیع پسین بر روی مشاهدات نمونه شرطی شده است و از اطلاعات و دانسته های مربوط به θ ساخته می شود، که به صورت یک کمیت و متغیر در نظر گرفته می شود. در دیدگاه بیزی تمام استنباط ها بر اساس توزیع پسین طراحی و ساخته می شوند. برای مثال، میانگین توزیع پسین می تواند به عنوان برآوردی برای پارامتر θ مورد استفاده قرار گیرد.

استنباط بیزی تقریبی

اغلب مسائل استنباطی در دیدگاه بیزی به محاسبه انتگرال هایی بر اساس توزیع پسین منتهی می شوند. فرض کنید محاسبه ی امیدریاضی توزیع پسین تابعی از پارامتر، یعنی

$$E[g(\theta)|z] = \int_{\theta} g(\theta)\pi(\theta|z) d\theta = \frac{\int_{\theta} g(\theta)f(z|\theta)\pi(\theta) d\theta}{\int_{\theta} f(z|\theta)\pi(\theta) d\theta} \quad (13.1)$$

^{۵۸}Posterior distribution

^{۵۹}Bayes' rule

^{۶۰}Informative

^{۶۱}Noninformative

^{۶۲}Vague

^{۶۳}Improper

مد نظر باشد. به طور مثال، برای محاسبه‌ی احتمال پسین یک مجموعه یا آزمون فرضیه، $g(\cdot)$ به صورت یک تابع نشان گر در نظر گرفته می شود. در مدل های پیچیده با توزیع های پسین پیچیده، معمولاً محاسبه رابطه‌ی (۱۳.۱) دشوار خواهد بود و نمی توان آن را به شکل بسته حساب کرد. اگر توزیع پیشین و توزیع پسین هر دو متعلق به یک خانواده باشند، اصطلاحاً این خانواده از توزیع ها را خانواده‌ی توزیع های مزدوج می نامند. انتخاب توزیع پیشین از خانواده توزیع های مزدوج می تواند منجر به محاسبات ساده برای توزیع پسین و رابطه (۱۳.۱) شود. اما در حالت کلی پیچیدگی محاسباتی یکی از مشکلات اساسی دیدگاه بیزی است که برای برخورد با آن از روش های تقریب توزیع پسین، مانند روش های شبیه سازی MCMC^{۶۴} و تقریب لاپلاس آشیانه‌ای جمع بسته^{۶۵} (INLA) استفاده می کنند (رو و همکاران، ۲۰۰۹؛ متروپولیس و همکاران، ۱۹۵۳؛ هستینگز، ۱۹۷۰).

۵.۱ پیشگویی فضایی

یکی از اهداف اصلی در تحلیل داده های فضایی، پیش گویی پاسخ در مکان های فاقد مشاهده است که با استفاده از روش کریگینگ انجام می شود (کرسبی، ۱۹۹۳). اصطلاح کریگینگ به افتخار دانیال کریگ^{۶۶} (۲۰۱۳ – ۱۹۱۹) نام گذاری شده است. کریگ در استفاده از روش های آماری برای ارزیابی و تحلیل مخازن معدنی پیشگام بود و بعدها این روش ها توسط ماترون^{۶۷} و دیگر همکارانش در دانشکده‌ی معدن پاریس فرمول بندی، توسعه و نشر داده شدند (دیگل و جورجی، ۲۰۱۹).

کریگینگ بهترین پیش گوی خطی نقطه‌ای و بلوکی است، به طوری که بهترین در این جا به معنی پیش گو با کمترین واریانس است. پیش گوی کریگینگ یک میانگین وزنی است که به نقاط نزدیکتر وزن بیشتر و به نقاط دورتر وزن کمتری اختصاص می دهد (اولیور و وستر، ۲۰۱۵). یکی از پذیره های روش کریگینگ، مانایی میدان تصادفی است. با معلوم بودن ساختار همبستگی فضایی و در نظر گرفتن تابع زیان توان دوم خطا، پیش گوی بهینه (کریگینگ) از می نیمم کردن میانگین توان دوم خطای پیش گو به دست می آید. در روش های کلاسیک، میدان تصادفی را گوسی در نظر می گیرند. بر حسب ساختار فضایی میدان تصادفی، سه نوع کریگینگ تعریف می شوند: ۱) کریگینگ ساده^{۶۸} که در آن فرض می شود میانگین میدان معلوم است؛ ۲) کریگینگ معمولی^{۶۹} که در آن فرض می شود میانگین ثابت ولی نامعلوم است؛ و ۳)

^{۶۴}Markov chain Monte Carlo

^{۶۵}Integrated nested Laplace approximation

^{۶۶}Daniel G. Krige

^{۶۷}Matheron

^{۶۸}Simple kriging

^{۶۹}Ordinary kriging

کریگینگ عام^{۷۰} که در آن میانگین به عنوان تابعی از موقعیت‌ها (و سایر متغیرهای تبیینی) در نظر گرفته می‌شود و به آن روند می‌گویند. در ادامه، به‌طور مختصر دو نوع کریگینگ معمولی و عام را معرفی می‌کنیم.

۱.۵.۱ کریگینگ معمولی

فرض کنید مقادیر متغیر پاسخ Z را در n مکان ناحیه تحت مطالعه، $D \subset \mathbb{R}^d$ ، مثل $\{s_1, \dots, s_n\}$ ، مشاهده کرده باشیم و هدف پیش‌گویی مقدار پاسخ در مکان جدید s_0 ، یعنی $Z(s_0)$ ، باشد. در کریگینگ معمولی تغییرات مشاهدات پاسخ از نظر فضایی به هم وابسته، تصادفی و تنها در یک مقدار نامعلوم μ ، مشترک هستند. در این حالت می‌توان میدان تصادفی گوسی $\{Z(s) : s \in D\}$ را به‌صورت زیر تجزیه کرد:

$$Z(s) = \mu + \delta(s) \quad (14.1)$$

که در آن μ ثابت و نامعلوم و $\delta(\cdot)$ میدان تصادفی مانای گوسی با میانگین صفر است. کریگینگ معمولی، به عنوان یک میانگین وزنی از مشاهدات، به‌صورت زیر تعریف می‌شود

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

ضرایب λ_i وزن‌های کریگینگ هستند که سهم هر کدام از مشاهدات را (بسته به فاصله آن‌ها تا مکان جدید s_0) در پیش‌گویی پاسخ یدک می‌کشند. چنان‌چه شرط $\sum_{i=1}^n \lambda_i = 1$ را در نظر بگیریم، کریگینگ یک پیش‌گوی نارایب خواهد بود، یعنی $E[\hat{Z}(s_0)] = \mu$. مقادیر بهینه ضرایب $\lambda = (\lambda_1, \dots, \lambda_n)^T$ از حل معادله

$$\lambda^T = (\gamma + \mathbf{1}_n^T \Gamma^{-1} \mathbf{1}_n)^{-1} (\mathbf{1}_n^T \Gamma^{-1} \gamma) \quad (15.1)$$

به‌دست می‌آیند، که در آن Γ یک ماتریس $n \times n$ با درایه (i, j)

$$\gamma(s_i - s_j) = \frac{1}{\tau} \text{Var}(Z(s_i) - Z(s_j))$$

و $\gamma = (\gamma(s_0 - s_1), \gamma(s_0 - s_2), \dots, \gamma(s_0 - s_n))^T$ هستند، به‌طوری که $\gamma(\cdot)$ تابع نیم‌تغییرنگار است. همچنین $\mathbf{1}_n$ برداری به طول n از ۱ها است. در مورد چگونگی رسیدن به این ضرایب و نحوه محاسبه واریانس کریگینگ معمولی می‌توانید به (کرس، ۱۹۹۳) مراجعه کنید.

کریگینگ ساده حالت خاصی از کریگینگ معمولی است که در آن فرض می‌شود، بین مشاهدات روندی وجود ندارد و میانگین معلوم و مستقل از مختصات فضایی است. در این حالت شرط نارایبی صفر بودن مجموع ضرایب لازم نیست و معادله محاسبه ضرایب کریگینگ، مشابه (۱۵.۱)، با کمی تغییرات جزئی، نتیجه می‌شود.

^{۷۰} Universal kriging

۲.۵.۱ کریگینگ عام

کریگینگ عام نسخه تعمیم یافته روش کریگینگ معمولی است. در واقعیت، در کاربردهای متعدد، مشاهدات یک میدان تصادفی دارای روند هستند و میانگین در موقعیت‌های مختلف متغیر است و این مهم باید لحاظ شود. کریگینگ عام به عنوان یک پیش گوی ناریب برای وقتی که روند میدان تصادفی یک ترکیب خطی نامعلوم از توابع معلوم است، استفاده می شود. در حالتی که میدان تصادفی $Z(s)$ دارای روند باشد، مدل به صورت

$$Z(s) = \sum_{j=1}^{p+1} \varphi_{j-1}(s) \beta_{j-1} + \delta(s), \quad s \in D \quad (۱۶.۱)$$

در نظر گرفته می شود، که در آن $\varphi_j(s)$ را تابع بنیادی^{۷۱} می نامند که بر حسب ماهیت روند در داده ها تعیین می شود. بردار ضرایب $\beta = (\beta_0, \dots, \beta_p)^T$ برداری نامعلوم از پارامترها، مجموعه ای معلوم از توابع $\delta(\cdot)$ یک میدان تصادفی مانای گوسی با میانگین صفر است. در این حالت، تابع روند $\mu(s)$ چند جمله ای

$$\mu(s) = \sum_{j=1}^{p+1} \varphi_{j-1}(s) \beta_{j-1}, \quad s \in D$$

است، با این فرض که $\varphi_0(s) = 1$ ، ضرایب λ به صورت

$$\lambda^T = \{\gamma + \mathbf{X}(\mathbf{X}^T \Gamma^{-1} \gamma)^{-1} (s - \mathbf{X}^T \Gamma^{-1} \gamma)\}^T \Gamma^{-1}$$

به دست می آیند، که در آن γ و Γ مشابه قبل تعریف می شوند و $\mathbf{X}$ ماتریسی $n \times (p+1)$ با درایه های

$$x_{ij} = \varphi_{j-1}(s_i), \quad j = 1, \dots, p+1, \quad i = 1, \dots, n$$

است (کرسی، ۱۹۹۳).

^{۷۱} Basis function

فصل ۲

مدل‌های سه بعدی زمین‌آماري

۱.۲ مقدمه

دنيایی که ما در آن زندگی می‌کنیم مملو از وابستگی‌هاست. در تحقیقات نادیده گرفتن این وابستگی می‌تواند نتایج غلط و زیان‌های گزافی را بر پژوهشگر و کاربران آن عرصه متحمل سازد. در علوم مختلف از جمله نفت، زمین‌شناسی، بوم‌شناسی گیاهی و حیوانی، سنجش از دور^۱ و پزشکی، اغلب داده‌هایی که استفاده می‌شوند، به موقعیت مکانی و جهت قرارگیری آن‌ها بستگی دارد. به‌طور معمول در اکثر موارد هر چه فاصله بین مشاهدات کمتر باشد شباهت بین آن‌ها بیشتر و در فاصله بیشتر شباهت کمتر است. تحلیل این نوع داده‌ها که به داده‌های فضایی شهرت دارند با روش‌های استنباطی آمار کلاسیک ناممکن است. وارد کردن ساختار وابستگی القاشده توسط موقعیت‌های مکانی داده‌ها (ساختار وابستگی فضایی) در تحلیل داده‌ها توسط روش‌های آمار فضایی ممکن شده است. یکی از زیرشاخه‌های آمار فضایی، زمین‌آمار است و به مدل‌ها و روش‌هایی اشاره دارد که از دو مشخصه پیروی می‌کنند. اول اینکه Z_i ها، $i = 1, \dots, n$ ، داده‌هایی هستند که در یک مجموعه گسسته مشتمل بر مکان‌های s_i در برخی از مناطق فضایی مانند D مشاهده شده‌اند. دوم هر مقدار مشاهده‌شده Z با اندازه‌گیری مستقیم یا با روابط آماری برآورد شده باشند (دیگل و ریبیرو، ۲۰۰۷). داده‌های زمین‌آماري در موقعیت‌های ثابت و در ناحیه‌ای پیوسته مشاهده می‌شوند و متغیر مورد بررسی می‌تواند پیوسته

^۱Remote Sensing

یا گسسته باشد. در حقیقت زمین‌آمار از ابتدای دهه ۱۹۵۰ به بعد در صنعت معدن توسعه یافت که در حال حاضر به‌طور گسترده‌ای در سایر علوم زیست محیطی برای نقشه‌برداری، نظارت و مدیریت استفاده می‌شود. همچنین با استفاده از مدل‌سازی همبستگی فضایی، زمین‌آمار می‌تواند پیش‌بینی‌های نااریب با کمترین واریانس را ارائه دهد (اولیور و وستر، ۲۰۱۵). روش‌های مدل‌سازی مبتنی بر زمین‌آمار نقش مهمی را در صنعت اکتشاف (معدن و نفت) ایفا می‌کند. چرا که در این صنایع با در نظر گرفتن هزینه‌های زیادی که صرف حفاری می‌شود، شناخت ناحیه مورد مطالعه و پیشگویی سه بعدی آن می‌تواند به کم کردن هزینه‌ها و کاهش خطا کمک کند. پیش‌بینی ویژگی‌های الاستیک ژئومکانیکی مخازن نفتی، به‌طور گسترده، نقش مهمی در بهینه‌سازی و کاهش خطرات و افزایش بهره‌وری حفاری در یک مخزن ایفا می‌کند (نظری استاد و همکاران، ۲۰۱۸). هر چاه دارای طول و عرض جغرافیایی است و از طرفی هزینه‌های حفاری با افزایش عمق به‌صورت نمایی رشد می‌کند (عبدآل و ال‌سهلاوی، ۲۰۱۴). مدل‌سازی سه بعدی متغیرهای تاثیرگذار در حفاری می‌تواند علاوه بر افزایش سرعت و کاهش هزینه، منجر به برنامه‌ریزی بهتر فرآیند حفاری و کاهش مخاطرات طبیعی شود. با این حال، ساخت مدل‌های سه بعدی قابل اعتماد برای این ویژگی‌ها، همچنان با سوال‌ها و گمانه‌هایی همراه است که پاسخ‌گویی به آن‌ها مستلزم تحقیق و توسعه مدل‌های کارا در میدان‌های تصادفی سه بعدی است.

در این فصل، ابتدا مسأله کاربردی مورد نظر پایان‌نامه در مورد مدل‌بندی سه بعدی مقاومت فشاری تک‌محوری سنگ در یکی از مخازن نفتی ایران را تشریح و هدف اصلی را بیان می‌کنیم. سپس تحلیل‌های موجود معمول در این حوزه را ارائه می‌دهیم.

۲.۲ تشریح مسأله

مطالعه رفتارهای سنگ مخزن با توجه به ویژگی‌های فیزیکی و مکانیکی آن، می‌تواند الگویی از تعادل سنگ را نشان دهد. دخالت‌های انسانی مانند حفاری، تولید یا تزریق ممکن است موجب برهم خوردن تعادل موجود در ساختار سنگ شود که نتایج این تغییرات را در حوزه مکانیک سنگ می‌توان مطالعه کرد (زوباک، ۲۰۰۷). استفاده از مکانیک سنگ برای محیط‌های زیرزمینی، ژئومکانیک نامیده می‌شود. اولین گام در هر فرآیند مربوط به ژئومکانیک، ایجاد یک مدل مکانیکی زمین^۲ (MEM) است. در واقع MEM یک نمایش عددی از خواص مکانیکی زمین است. مدل‌سازی سه بعدی ژئومکانیکی به بررسی پارامترهای ژئومکانیکی مخزن در سه بعد می‌پردازد و از این پارامترها در مقیاس حجمی استفاده می‌شود. پیش‌بینی ویژگی‌های الاستیک ژئومکانیکی مخازن نفتی، به‌طور گسترده، نقش مهمی در بهینه‌سازی و کاهش خطرات و افزایش بهره‌وری حفاری در یک مخزن ایفا می‌کند (نظری استاد و همکاران، ۲۰۱۸).

^۲Mechanical earth model

با این حال، ساخت مدل‌های سه بعدی قابل اعتماد برای این ویژگی‌ها همچنان با سوال‌ها و گمانه‌هایی همراه است که پاسخ‌گویی به آن‌ها مستلزم تحقیق و توسعه مدل‌های کارا است. روش‌های متنوعی برای مدل‌سازی ژئومکانیک وجود دارد که از جمله می‌توان به زمین‌آمار اشاره کرد.

داده‌های استفاده شده در این پژوهش مربوط به مقاومت فشاری تک محوری سنگ^۳ (UCS) است. مقاومت فشاری نهایی سنگ، مقداری است که تنش‌های فشاری تک‌محوری سبب شکست کامل سنگ می‌شود. یکی از پارامترهای مهم در ارزیابی مقاومت سنگ و ساخت مدل ژئومکانیک، مقاومت تک محوری سنگ بکر UCS است. معمولاً مقدار مقاومت فشاری از طریق آزمایش به‌دست می‌آید. از طرفی دسترسی به مغزه حفاری مشکل و هزینه‌بر است، به همین دلیل محققان صنعت نفت روابطی را برای برآورد پارامترهای مکانیکی از روی نگاره‌های پتروفیزیکی ارائه داده‌اند. در این پژوهش برای محاسبه مقدار مقاومت فشاری تک محوری از رابطه‌ی پیشنهادی توسط شرکت ملی مناطق نفت‌خیز جنوب ایران که برای مخازن نفتی کشور ارائه شده، استفاده کردیم. مقدار UCS به صورت

$$UCS = 2/27E_{sta} + 4/74$$

محاسبه می‌شود که در آن E_{sta} مدول یانگ استاتیک است و از رابطه‌ی $E_{sta} = 0.7E_{Dyn}$ به‌دست می‌آید و E_{Dyn} مدول یانگ دینامیک است و از رابطه‌ی

$$E_{Dyn} = \frac{\rho V_s^2 (3V_p^2 - 4V_s^2)}{3V_p^2 - 4V_s^2}$$

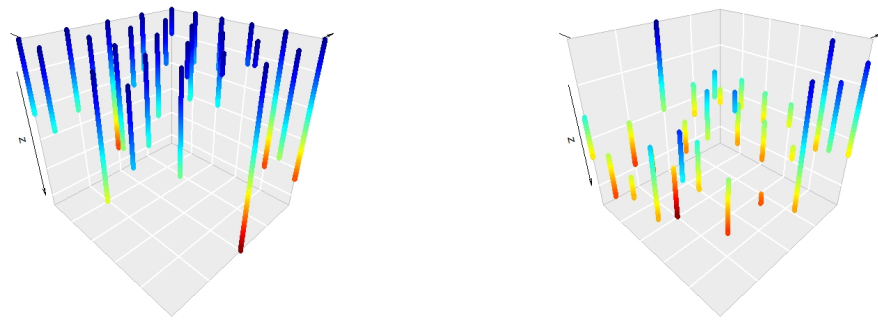
به‌دست می‌آید، که در آن V_p^2 و V_s^2 به ترتیب سرعت موج برشی و فشاری بر حسب $\frac{Km}{sec}$ و ρ دانسیته سنگ بر حسب $\frac{gm}{cc}$ است.

منطقه مورد مطالعه در این پایان‌نامه، یکی از مخازن نفتی ایران شامل ۲۷ حلقه چاه با طول و عرض جغرافیایی و عمق‌های مشخص است که مقدار UCS با استفاده از متغیرهای ثبت‌شده‌ی مرتبط با این متغیر به‌طور تقریبی محاسبه شده است. لازم به ذکر است که چاه‌ها دارای ضخامت‌های متفاوتی هستند. از طرفی این مخزن بر روی یک چین زمین‌شناسی قرار دارد، بنابراین از رابطه‌ی (۱.۲) برای یکسان‌سازی عمق چاه‌ها در همان مختصات طول و عرض جغرافیایی استفاده کردیم.

$$s_{(new)i} = s_{(old)i} - |\min(s_{ij}) - \min(s_{(old)i})|, \quad j = 1 \dots 10, \quad i = 1, \dots, m_j \quad (1.2)$$

که $s_{(old)i}$ و $s_{(new)i}$ به ترتیب مختصات عمق حقیقی و اصلاح شده برای مکان i ام در هر چاه است، همچنین $\min(s_{ij})$ کمترین عمق اندازه‌گیری‌شده در تمام چاه‌ها و m_j نشان‌دهنده‌ی تعداد مشاهدات در هر چاه است. در شکل ۱.۲ نمودار سه بعدی چاه‌ها با عمق حقیقی و عمق اصلاحی نمایش داده شده‌اند. از بین چاه‌ها پنج چاه که شامل ۱۸۱۴۰ مشاهده است را

^۳Unconfined compressive strength



(آ) نمودار سه بعدی چاه‌ها با عمق حقیقی (ب) نمودار سه بعدی چاه‌ها بعد از یکسان‌سازی عمق

شکل ۱.۲: نمودار سه بعدی چاه‌ها

به‌طور تصادفی برای مجموعه داده‌های آزمون^۴ انتخاب کردیم تا نتایج را بر اساس اعتبارسنجی متقابل^۵ مورد ارزیابی قرار دهیم و از بقیه چاه‌ها که ۵۸۴۶۱ مشاهده را شامل می‌شود، به‌عنوان مجموعه آموزشی^۶ برای مدل‌سازی سه بعدی استفاده کردیم. شکل ۲.۲ مربوط به محل قرارگیری چاه‌ها در موقعیت‌های مکانی طول و عرض جغرافیایی به تفکیک مجموعه‌های آموزشی و آزمون است.

۳.۲ ساختار همبستگی فضایی

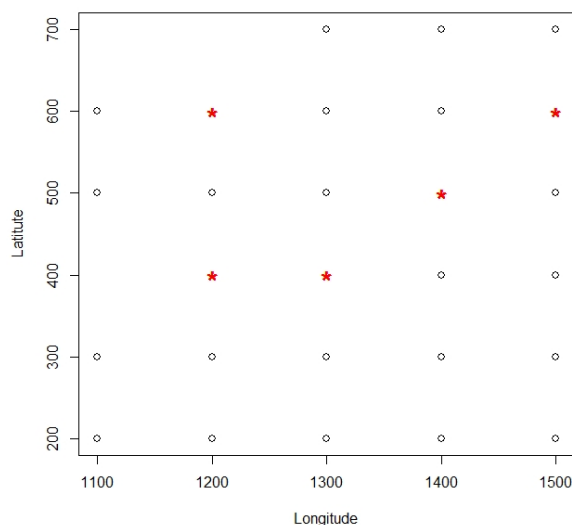
همان‌طور که در فصل اول بیان کردیم، در آمار فضایی از تغییرنگار و هم‌تغییرنگار برای بررسی پراکندگی و شباهت داده‌ها استفاده می‌شود و می‌توان گفت که مفاهیمی شبیه به واریانس و کوواریانس در آمار کلاسیک دارند. در زمین‌آمار توصیف نحوه‌ی تغییرپذیری مکانی داده‌ها با استفاده از تابع تغییرنگار را واریوگرافی^۷ می‌نامند که شامل رسم تغییرنگار تجربی، برازش مدل مناسب به آن و تحلیل تغییرنگار مشخص‌شده برای ساختار تغییرات داده‌ها است. اهمیت واریوگرافی به این دلیل است که تمامی تحلیل‌ها و استنباط‌های آماری از جمله پیشگویی، با تکیه بر پارامترهای حاصل از این عملیات انجام می‌شوند.

^۴Test set

^۵Cross validation

^۶Train set

^۷Variography



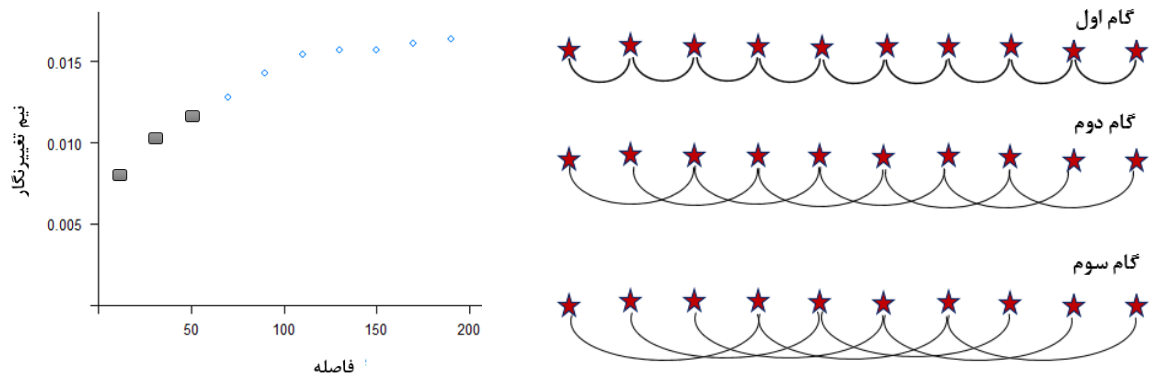
شکل ۲.۲: محل قرارگیری چاه‌های مجموعه آموزشی (توخالی) و مجموعه آزمون (ستاره‌ای) بر حسب طول و عرض جغرافیایی و محل قرارگیری آن‌ها

۱.۳.۲ برآورد تغییرنگار مخزن نفتی

برآورد دقیق تغییرنگار برای پیش‌گویی قابل اعتماد توسط کریگینگ مورد نیاز است. قبلاً ذکر شده که تغییرنگار تجربی (نمونه) معمولاً توسط روش‌های گشتاوری در دنباله‌ای از گام‌ها محاسبه می‌شود و یک یا چند مدل معتبر تغییرنگار بر روی آن برازش داده می‌شود. با توجه به ماهیت سه بعدی میدان تصادفی مخزن نفتی و اهمیت برآورد ساختار وابستگی مقاومت فشاری تک‌محوری توسط تغییرنگار، روند برآورد تغییرنگار را از فضای یک بعدی تا سه بعدی تشریح می‌کنیم.

برآورد تغییرنگار یک بعدی

نمونه‌گیری در یک بعد ممکن است به صورت برش‌های عمودی یا افقی باشد. گام h به اِزاء $h = |h|$ مقادیر عددی را ارائه می‌دهد و جایگزین مقدار h در معادله‌ی (۵.۱) می‌شود. نیم‌تغییرنگار، $\hat{\gamma}(h)$ ، می‌تواند فقط به صورت چندگانه در بازه‌ی نمونه محاسبه شود. شکل ۳.۲ (آ) به صورت مقایسه‌ای، سه گام $h = 1, 2, 3$ بین جفت نقاط را نشان می‌دهد. نتایج به صورت مجموعه‌ای از نیم‌تغییرنگارهای $\hat{\gamma}(1), \hat{\gamma}(2), \hat{\gamma}(3), \dots$ است. نمودار نیم‌تغییرنگار یک بعدی تجربی ده گام اول در شکل ۳.۲ (ب) آورده شده است.



(ب) نمودار نیم‌تغییرنگار تجربی برای ۱۰ گام اول (آ) نمایش سه گام $h = 1, 2, 3$ بین جفت نقاط

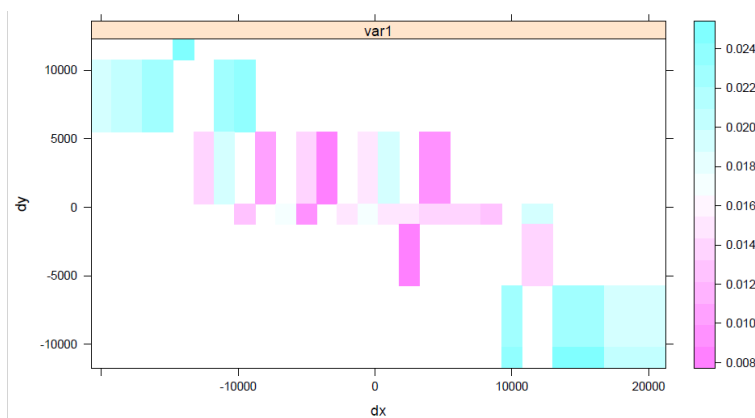
شکل ۳.۲: تغییرنگار یک بعدی

برآورد تغییرنگار دو بعدی

برآورد تغییرنگار در دو بعد طی سه مرحله انجام می‌شود. ابتدا باید ناحیه را به شبکه‌های دو بعدی تقسیم کنیم. این کار می‌تواند موجب تشخیص همسانگردی میدان با استفاده از چرخش در جهت‌های مختلف تغییرات شود. می‌توان تغییرنگار را در شبکه‌ها در جهت‌های مختلف محاسبه کرد. به عنوان مثال در جهت طول، عرض یا در جهت قطر. دوم، تغییرنگار دو بعدی به صورت زیر محاسبه می‌شود:

$$\begin{aligned} \hat{\gamma}(p, q) &= \frac{1}{2(m-p)(n-q)} \sum_{i=1}^{m-p} \sum_{j=1}^{n-q} \{z(i, j) - z(i+p, j+q)\}^2 \\ \hat{\gamma}(p, -q) &= \frac{1}{2(m-p)(n-q)} \sum_{i=1}^{m-p} \sum_{j=q+1}^n \{z(i, j) - z(i+p, j-q)\}^2. \end{aligned} \quad (2.2)$$

در این جا p و q به ترتیب تعداد گام‌های سطر و ستون در طول توری هستند. به طور کلی شبکه برابر با فاصله‌ی گام است. محاسبات تغییرنگار از q تا $-q$ و از p تا p است. نقشه‌ی دو بعدی تغییرنگار در شکل ۴.۲ نشان داده شده است. سوم، تغییرنگار می‌تواند در تمام جهت‌ها برای شبکه‌های منظم یا نامنظم محاسبه شود.



شکل ۴.۲: اجزاء تشکیل دهنده تغییرنگار

برآورد تغییرنگار سه بعدی

محاسبه تغییرنگار سه بعدی در میدان با توری‌های حجمی منظم یا نامنظم به این شکل است که می‌توان توری را به صورت یک سری در نظر گرفت و در سه بعد طول، عرض و عمق مورد بررسی قرار داد. معادله‌ی (۳.۲) مربوط به تغییرنگار سه بعدی است که برای بررسی همسانگردی باید در چند جهت مختلف از طول، عرض و عمق یک توری محاسبه شود.

$$\hat{\gamma}(p, q, r) = \frac{1}{\sqrt{(m-p)(n-q)(l-r)}} \sum_{i=1}^{m-p} \sum_{j=1}^{n-q} \sum_{k=1}^{l-r} \{z(i, j, k) - z(i+p, j+q, k+r)\}^2 \quad (3.2)$$

که در آن p, q, r به ترتیب اندازه گام‌های طول، عرض و عمق و m, n, l تعداد گام‌های طول، عرض و عمق هستند.

۲.۳.۲ برآوردگر استوار

برآوردگر گشتاوری که در معادله‌ی (۵.۱) آورده شده است وقتی داده‌ها فاقد روند باشند، نارایب است، یعنی $E[\hat{\gamma}] = \gamma$. اما این برآوردگر استوار^۸ نیست و دارای خواص پایداری ضعیفی است. چرا که برای فاصله‌های بزرگ از داده‌های کمتری استفاده می‌کند و نیز نسبت به نقاط پرت و دورافتاده بسیار تاثیرپذیر است. **کرسی و هاوکینز (۱۹۸۰)** دو برآوردگر استوار را معرفی کردند که از معادلات

$$\hat{\gamma}(h) = \frac{\left\{ \frac{1}{N} \sum_{N(h)} |Z(s_i) - Z(s_j)|^{\frac{1}{\sqrt{2}}} \right\}^4}{0.914 + \frac{0.988}{N(h)}} \quad (4.2)$$

و

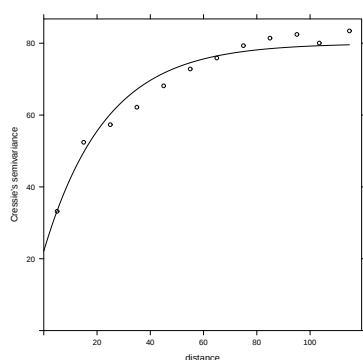
$$\hat{\gamma}(h) = \frac{(Med\{|Z(s_i) - Z(s_j)|^{\frac{1}{\sqrt{2}}} : (s_i, s_j) \in N(h)\})^4}{0.914} \quad (5.2)$$

^۸Robust

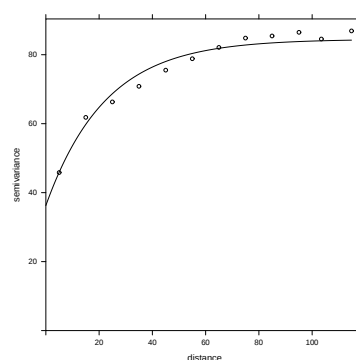
پیروی می‌کند، که در آن $Med(\cdot)$ مقدار میانه در مجموعه $N(h)$ است.

نتایج برآورد تغییرنگار سه بعدی برای میدان نفتی

برازش مدل مناسب به تغییرنگار امر مهمی است که در پیشگویی تاثیر بسزایی دارد. در این جا از دو روش رسم نمودار (روش چشمی) و مقایسه مجموع توان‌های دوم خطا^۹ (SSErr) برای انتخاب مدل مناسب تغییرنگار تجربی استفاده کردیم. بعد از رسم نمودار نیم‌تغییرنگار تجربی کلاسیک منحنی مدل‌هایی که به نظر مناسب بودند به نمودار نیم‌تغییرنگار برازش دادیم. سپس مدلی که دارای کمترین مقدار SSErr بود را انتخاب کردیم. در این جا مدل پارامتری نمایی، SSErr کمتری نسبت به سایر مدل‌ها داشت. برای اطمینان بیشتر از برآوردگر استوار هم استفاده کردیم، که مقدار SSErr بیشتری نسبت به برآوردگر کلاسیک داشت. نمودار ۵.۲ منحنی برآوردگر کلاسیک و استوار مدل پارامتری نمایی برازش داده شده به نیم‌تغییرنگار را نشان داده‌است. مقدار ازاره، دامنه و اثر قطعه‌ای به ترتیب برابر با ۸۴/۵، ۶۷/۲۳ و ۳۶/۲۷ است.



(ب)



(آ)

شکل ۵.۲: نمودار سه بعدی نیم‌تغییرنگار تجربی منحنی مدل پارامتری نمایی برازش داده شده (آ) برآوردگر کلاسیک، (ب) برآوردگر استوار

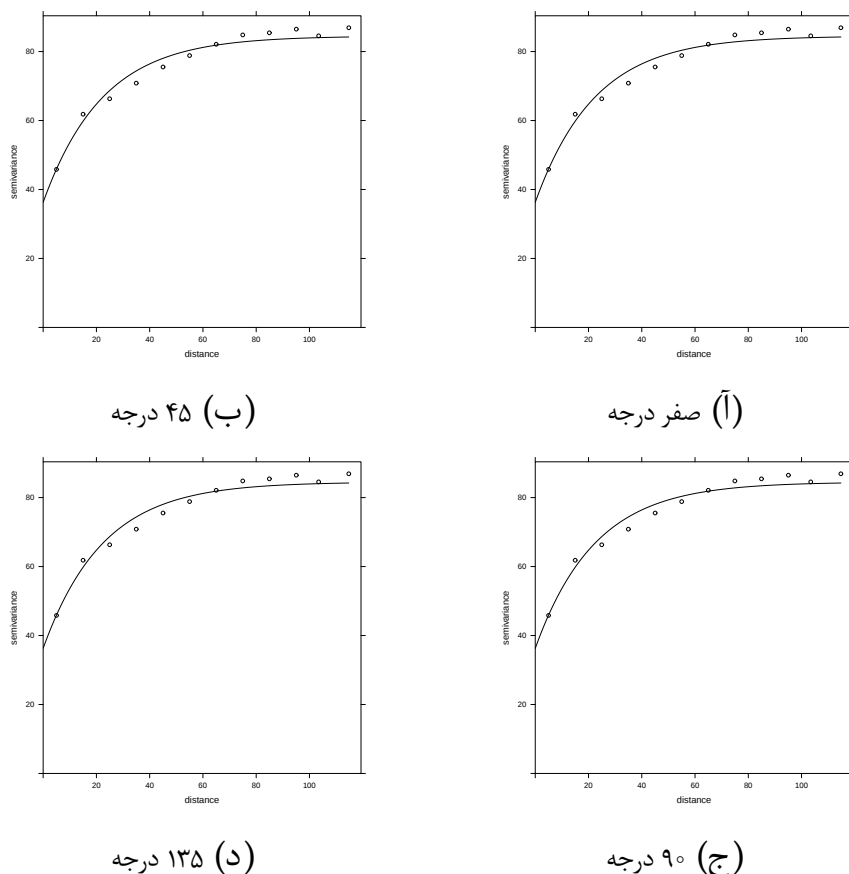
۳.۳.۲ بررسی ناهمسانگردی میدان نفتی

با توجه به اهمیت همسانگرد بودن میدان و تاثیر آن بر برآورد درست ساختار وابستگی، وجود ناهمسانگردی در میدان نفتی را بررسی می‌کنیم.

نمودار سه بعدی نیم‌تغییرنگار تجربی داده‌های مقاومت فشاری در چهار جهت ۰، ۴۵، ۹۰ و ۱۳۵ درجه برای بررسی ناهمسانگردی همراه با منحنی مدل پارامتری نمایی برازش داده‌شده

^۹Sum of squares error

در شکل ۶.۲ آورده شده است. با توجه به شکل می توان دریافت که نیم تغییرنگار در این چهار جهت مشابه و در نتیجه میدان همسانگرد است.

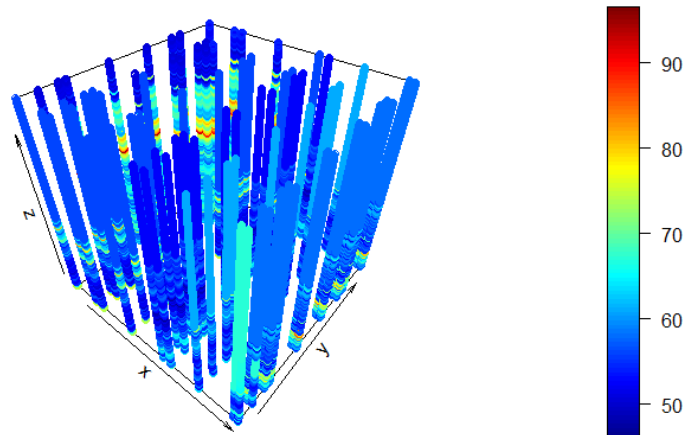


شکل ۶.۲: نمودار سه بعدی نیم تغییرنگار و منحنی مدل پارامتری نمایی برازش داده شده در چهار جهت

۴.۲ پیشگویی سه بعدی مقاومت فشاری تک محوری سازند

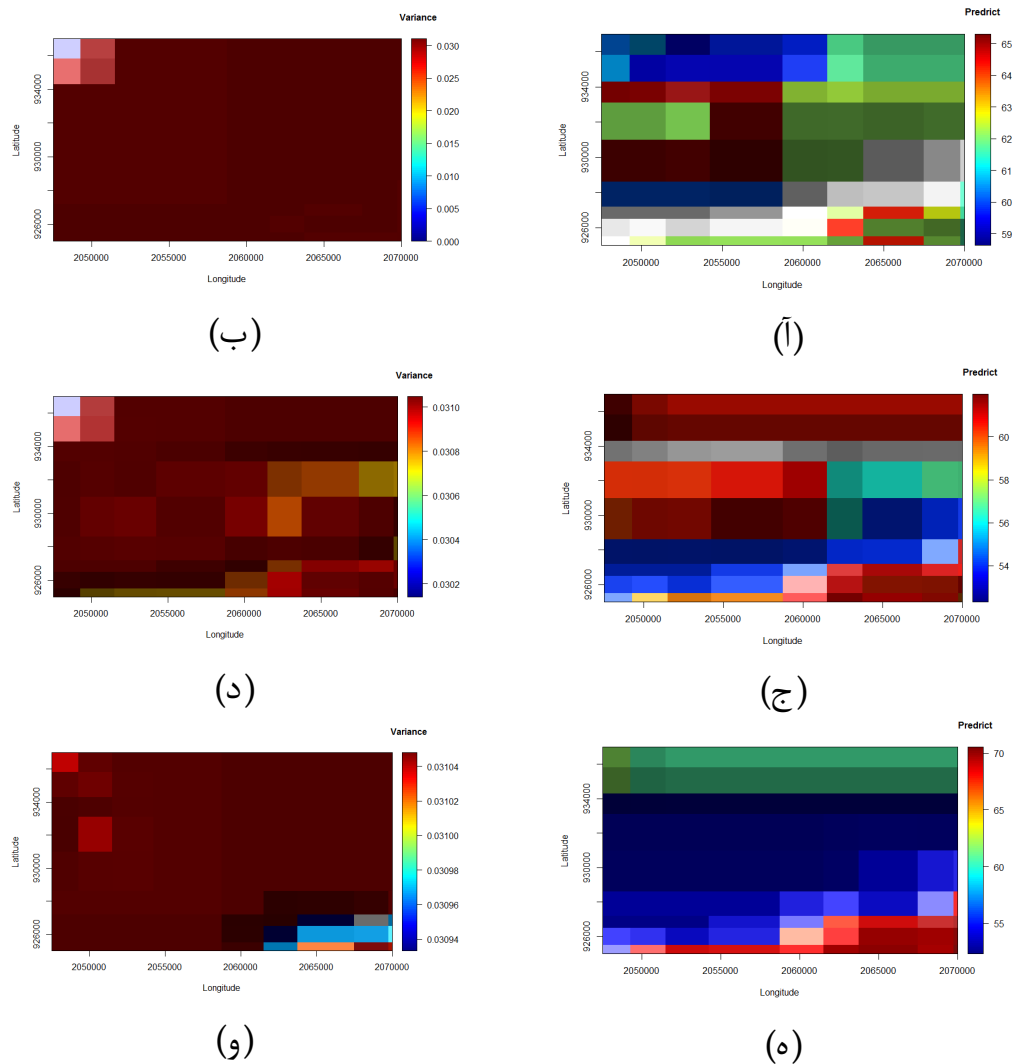
برای پیش گویی داده های مقاومت فشاری تک محوری سازند، با توجه به پذیره گوسی بودن میدان تصادفی (۱۴.۱)، ابتدا ناحیه تحت مطالعه را شبکه بندی کردیم. به این صورت که طول و عرض میدان نفتی را به ۶ قسمت مساوی تقسیم کردیم. سپس عمق ها را از کمترین مقدار تا بیشترین مقدارش به برش های ۱۵ سانتی متری افراز کردیم. از ضرب طول، عرض و عمق یک شبکه که به مکعب های مستطیلی کوچکی به طول ۱۷۴، عرض ۸۳ و ارتفاع ۱۵٪ متر افراز شده است به دست می آید. شبکه ۹۵۰۰۰ گره را شامل می شود. سپس مختصات جغرافیایی ۵ چاهی که به عنوان داده های آزمون در نظر گرفتیم را به این شبکه برای پیشگویی

اضافه کردیم. در شکل ۷.۲ با استفاده از کریگینگ معمولی و شبکه‌بندی ناحیه تحت مطالعه،

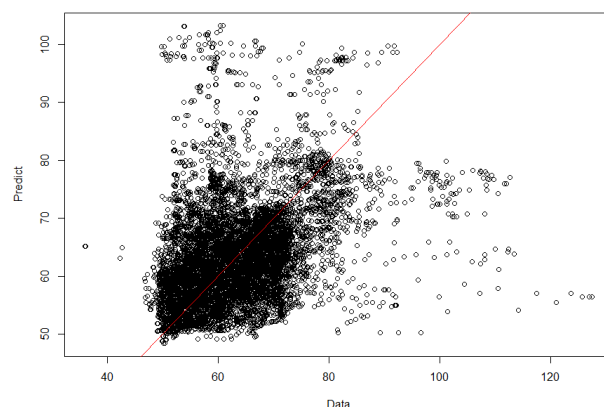


شکل ۷.۲: پیش‌گویی سه بعدی میدان نفتی تحت مطالعه

نقشه پهنه‌بندی مقادیر پیش‌گویی مقدار مقاومت فشاری را به‌دست آوردیم. شکل ۸.۲ نمودار دو بعدی مقادیر پیش‌گویی و واریانس پیش‌گویی را در سه عمق مختلف نشان می‌دهد. محور افقی شکل ۹.۲ نشان دهنده‌ی مقادیر واقعی مجموعه داده‌های آزمون Z_i و محور عمودی مقادیر پیش‌گویی به دست‌آمده از کریگینگ معمولی $\hat{Z}_i$ برای مجموعه‌ی آزمون است. برای UCS کمتر از 8° توزیع پراکنش داده‌ها تقریباً حول نیم‌ساز ربع اول و سوم ($Z_i = \hat{Z}_i$) است و بیانگر مناسب بودن مدل و مقدار ناچیز باقی‌مانده‌ها است. با افزایش مقادیر داده‌ها، پیش‌گو و مقادیر اصلی از نیم‌ساز فاصله گرفته و اندازه‌ی خطا بیشتر شده است. به عبارتی واریانس خطا ثابت نیست ولی همچنان نیم‌ساز بین داده‌ها قرار دارد و عملکرد روش کریگینگ سه بعدی مناسب است.



شکل ۸.۲: نقشه پهنه بندی پیشگویی مقادیر مقاومت فشاری در عمق‌های ۱۷۰۳، ۲۱۵۵ و ۲۴۸۰ متر به ترتیب در (آ)، (ج) و (ه). همچنین نقشه پهنه‌بندی واریانس پیشگویی در همین عمق‌ها به ترتیب در (ب)، (د) و (و).



شکل ۹.۲: مقادیر پیشگویی در مقابل مقادیر واقعی داده‌های آزمون

فصل ۳

روش‌های ناپارامتری فضایی در میدان‌های تصادفی سه بعدی

در این فصل کریگینگ میانه پالایش سه بعدی را توضیح می‌دهیم. سپس از دو مثال شبیه‌سازی شده برای پیش‌گویی با استفاده از کریگینگ عام که در آن از مدل‌بندی اسپلاین استفاده شده و کریگینگ میانه پالایش بهره‌می‌بریم. در نهایت به مقایسه‌ی روش‌های مذکور می‌پردازیم.

۱.۳ میانه پالایش

توکی (۱۹۷۰) و **ولمن و هوگلین (۱۹۸۱)** میانه پالایش دو طرفه^۱ را روشی برای استخراج اثرات ستونی و سطری با به‌کارگیری میانه به‌جای میانگین معرفی کردند. میانه پالایش کاربردهای زیادی دارد. برای نمونه در آمار فضایی، کرسی در سال ۱۹۹۳ با استفاده از شبکه‌بندی داده‌ها از این روش به‌عنوان یک روش تنومند برای حذف هرگونه روند در دامنه فضایی استفاده کرد. میانه پالایش معمولاً برازشی را برای یک مدل جمعی ارائه می‌دهد که از دید کمترین قدرمطلق باقی‌مانده‌ها تقریباً بهینه است (**موستلر و توکی، ۱۹۷۷**). در این پایان‌نامه، برای کاربردهای سه بعدی در زمین‌شناسی و نفت، مشابه الگوریتم ارائه شده توسط **کرسی (۱۹۹۳)**، برای استخراج اثرات طول، عرض و بعد سوم (عمق یا ارتفاع) در تابع روند میدان تصادفی برای اولین بار یک

^۱Two-way

الگوریتم میانه‌پالایش سه بعدی را پیشنهاد و معرفی می‌کنیم و کارایی آن را با دو روش دیگر در کیفیت پیشگویی فضایی مقایسه می‌کنیم.

۱.۱.۳ الگوریتم میانه‌پالایش سه بعدی

در این بخش، روش کار را به صورت جبری توضیح می‌دهیم. برای درک بیشتر شماره‌ی هر مرحله از تکرار میانه‌پالایش را بالای نماد و داخل پرانتز قرار می‌دهیم. پس $R_a^{(n)}$ ، $C_b^{(n)}$ و $L_c^{(n)}$ نشان‌دهنده‌ی برآورد اثرات میانه در سطوح مختلف به ترتیب از عوامل سطر، ستون و لایه هستند که در مرحله‌ی n تأثیرشان بر روی داده‌ها حذف شده است. مقدار اولیه میانه را برای هر سه عامل برابر صفر قرار می‌دهیم. یعنی برای تمام a ، b و c ها

$$R_a^{(0)} = C_b^{(0)} = L_c^{(0)} = 0$$

اگر $r_{abc}^{(n)}$ باقی‌مانده‌ها پس از حذف اثرات سطری، ستونی و لایه‌ای در مرحله‌ی n باشند، پس می‌توانیم از $r_{abc}^{(0)} = Z(s_{abc})$ شروع کنیم که $Z(s_{abc})$ همان داده‌های اصلی در موقعیت‌های s_{abc} هستند. اکنون جزئیات مراحل میانه‌پالایش را مورد بررسی قرار می‌دهیم. میانه‌ها را باید برای هر عامل به ترتیب حذف کنیم. گاهی نتایج به ترتیب انتخاب عامل‌ها بستگی دارد. وقتی $i = 1, \dots, n$ را شمارنده‌ی هر تکرار در نظر می‌گیریم، هر مرحله از تکرار میانه‌پالایش از دستورات زیر تبعیت می‌کند: ابتدا میانه سطری را در هر سطح محاسبه می‌کنیم، یعنی

$$R_a^{(i)} = \text{med}_{bc} \{r_{abc}^{(i-1)}\} \quad a = 1, \dots, U$$

که در آن U تعداد کل سطرهاست. میانه را روی b و c حساب می‌کنیم و a را ثابت در نظر می‌گیریم. سپس میانه‌ها را از $r_{abc}^{(i-1)}$ کم می‌کنیم و میانه ستونی را حساب می‌کنیم.

$$C_b^{(i)} = \text{med}_{ac} \{r_{abc}^{(i-1)} - R_a^{(i)}\} \quad b = 1, \dots, V$$

که در آن V تعداد کل ستون‌هاست. در نهایت مقدار $C_b^{(i)}$ را کم می‌کنیم و میانه را برای هر مرحله از عامل لایه به تعداد کل W ، محاسبه می‌کنیم. داریم

$$L_c^{(i)} = \text{med}_{ab} \{r_{abc}^{(i-1)} - R_a^{(i)} - C_b^{(i)}\} \quad c = 1, \dots, W$$

با کم کردن مقدار $L_c^{(i)}$ باقی‌مانده‌ها در مرحله اول به صورت زیر به دست می‌آیند:

$$r_{abc}^{(i)} = r_{abc}^{(i-1)} - R_a^{(i)} - C_b^{(i)} - L_c^{(i)}.$$

اگر بخواهیم برآورد اثرات را در این مرحله مرکزی کنیم، میانه را برای هر مجموعه از اثر حساب

می‌کنیم. سپس برآورد اثرات به صورت زیر محاسبه می‌شوند:

$$\begin{aligned}\hat{R}_a &= \sum_{i=1}^n R_n^{(i)} - \text{med}_a \left\{ \sum_{i=1}^n R_a^{(i)} \right\}, \quad a = 1, \dots, U \\ \hat{C}_b &= \sum_{i=1}^n R_n^{(i)} - \text{med}_b \left\{ \sum_{i=1}^n C_b^{(i)} \right\}, \quad b = 1, \dots, V \\ \hat{L}_c &= \sum_{i=1}^n R_n^{(i)} - \text{med}_c \left\{ \sum_{i=1}^n L_c^{(i)} \right\}, \quad c = 1, \dots, W \\ \hat{\mu} &= \text{med}_a \left\{ \sum_{i=1}^n R_a^{(i)} \right\} + \text{med}_b \left\{ \sum_{i=1}^n C_b^{(i)} \right\} + \text{med}_c \left\{ \sum_{i=1}^n L_c^{(i)} \right\}.\end{aligned}$$

در این جا $\hat{\mu}$ برآورد میانه کل است. برآورد نهایی به ترتیب انتخاب عامل‌ها برای حذف اثر آن‌ها بستگی دارد. اما در کاربرد معمولاً این موضوع از اهمیت زیادی برخوردار نیست. در عمل برای روش میانه پالایش تکرارهای ثابتی در نظر گرفته می‌شود (حداقل ۲ یا ۳ تکرار) و از آن به عنوان تحلیل استاندارد استفاده می‌شود. در روش میانه پالایش سه طرفه^۲ محاسبات برای رسیدن به هدف شرط جانبی یا حذف کامل اثرات هدایت می‌شوند. یعنی برای رسیدن به وضعیتی که

$$\begin{aligned}\text{med}_{bc} \{r_{abc}\} &= \circ, \text{ برای } a \text{ هر} & \text{med}_a \{\hat{R}_a\} &= \circ \\ \text{med}_{ac} \{r_{abc}\} &= \circ, \text{ برای } b \text{ هر} & \text{med}_b \{\hat{C}_b\} &= \circ \\ \text{med}_{ab} \{r_{abc}\} &= \circ, \text{ برای } c \text{ هر} & \text{med}_c \{\hat{L}_c\} &= \circ\end{aligned} \tag{۱.۳}$$

روابط (۱.۳) نشان می‌دهند که اثرات سطری، ستونی و لایه‌ای در داده‌ها به طور کامل حذف شده‌اند.

مثال ۱.۱.۳. برای درک بیشتر الگوریتم میانه پالایش سه طرفه، از یک مثال ساده با داده‌های مصنوعی، استفاده می‌کنیم. ماتریس سه بعدی $(3 \times 2 \times 5)$ در (۲.۳) دارای ۳ سطر، ۲ ستون

^۲Three-way

و ۵ لایه است.

$$(2.3) \quad \begin{matrix} & a_1 & a_2 & a_3 \\ b_1 & \begin{pmatrix} c_1 & 3 \\ c_2 & 1 \\ c_3 & 6 \\ c_4 & 12 \\ c_5 & 9 \end{pmatrix} & \begin{pmatrix} -7 \\ -5 \\ 2 \\ 0 \\ 1 \end{pmatrix} & \begin{pmatrix} 5 \\ 8 \\ 2 \\ -1 \\ 0 \end{pmatrix} \\ b_2 & \begin{pmatrix} c_1 & 0 \\ c_2 & -5 \\ c_3 & 4 \\ c_4 & 12 \\ c_5 & 1 \end{pmatrix} & \begin{pmatrix} 3 \\ -1 \\ -10 \\ 6 \\ -8 \end{pmatrix} & \begin{pmatrix} -4 \\ -6 \\ 3 \\ 11 \\ 7 \end{pmatrix} \end{matrix}$$

با انتخاب سطر به عنوان اولین عامل تاثیرگذار بر روی داده‌ها، ابتدا میانه‌ی سطری را با ثابت در نظر گرفتن سطوح آن، روی سطوح ستون و لایه محاسبه می‌کنیم. میانه سطری متغیر^۳ را با نماد CRE و اثر سطری را با RE نشان می‌دهیم. مقدار اولیه اثرات سطری برای تمام سطوح را برابر با صفر قرار می‌دهیم.

$$\begin{aligned} CRE_1^{(1)} &= \text{median}(\{3, 1, 6, 12, 9\}, \{-7, -5, 2, 0, 1\}, \{5, 8, 2, -1, 0\}) = 2 \\ CRE_2^{(1)} &= \text{median}(\{0, -5, 4, 12, 1\}, \{3, -1, -10, 6, -8\}, \{-4, -6, 3, 11, 7\}) = 1 \\ RE_1^{(0)} &= RE_2^{(0)} = 0 \\ RE_1^{(1)} &= CRE_1^{(1)} - RE_1^{(0)} = 2 - 0 = 2 \\ RE_2^{(1)} &= CRE_2^{(1)} - RE_2^{(0)} = 1 - 0 = 1 \end{aligned}$$

میانه اثر سطری مرحله اول که با $RE^{(1)}$ نمایش می‌دهیم، به صورت زیر حساب می‌شود:

$$RE^{(1)} = \text{median}(RE_1^{(1)}, RE_2^{(1)}) = \text{median}(2, 1) = 1/5$$

سپس اثرات سطری را از داده‌های مربوط به همان سطر کم می‌کنیم که ماتریس (۳.۳) نتیجه می‌شود.

^۳Change row effect

$$\begin{matrix}
 & a_1 & a_2 & a_3 \\
 b_1 & \begin{pmatrix} c_1 & 1 \\ c_2 & -1 \\ c_3 & 4 \\ c_4 & 10 \\ c_5 & 7 \end{pmatrix} & \begin{pmatrix} -9 \\ -7 \\ 0 \\ -2 \\ -1 \end{pmatrix} & \begin{pmatrix} 3 \\ 6 \\ 0 \\ -3 \\ -2 \end{pmatrix} \\
 b_2 & \begin{pmatrix} c_1 & -1 \\ c_2 & -6 \\ c_3 & 3 \\ c_4 & 11 \\ c_5 & 0 \end{pmatrix} & \begin{pmatrix} 2 \\ -2 \\ -11 \\ 5 \\ -9 \end{pmatrix} & \begin{pmatrix} -5 \\ -7 \\ 2 \\ 10 \\ 6 \end{pmatrix}
 \end{matrix} \quad (3.3)$$

برای حذف اثر ستون مشابه سطر رفتار می‌کنیم، مقدار اولیه اثرات ستون^۴ (CE)، برای تمام سطوح برابر صفر است. اثر متغیر ستون^۵ را با CCE نمایش می‌دهیم.

$$\begin{aligned}
 CCE_1^{(1)} &= \text{median}(\{1, -1, 4, 10, 7\}, \{-1, -6, 3, 11, 0\}) = 2 \\
 CCE_2^{(1)} &= \text{median}(\{-9, -7, 0, -2, -1\}, \{2, -2, -11, 5, -9\}) = -2 \\
 CCE_3^{(1)} &= \text{median}(\{3, 6, 0, -3, -2\}, \{-5, -7, 2, 10, 6\}) = 1 \\
 CE_1^{(0)} &= CE_2^{(0)} = CE_3^{(0)} = 0 \\
 CE_1^{(1)} &= CCE_1^{(1)} - CE_1^{(0)} = 2 - 0 = 2 \\
 CE_2^{(1)} &= CCE_2^{(1)} - CE_2^{(0)} = -2 - 0 = -2 \\
 CE_3^{(1)} &= CCE_3^{(1)} - CE_3^{(0)} = 1 - 0 = 1
 \end{aligned}$$

اکنون میانه اثر ستونی مرحله اول که با $CE^{(1)}$ نشان می‌دهیم، به صورت زیر حساب می‌شود:

$$CE^{(1)} = \text{median}(CE_1^{(1)}, CE_2^{(1)}, CE_3^{(1)}) = \text{median}(2, -2, 1) = 1.$$

سپس اثرات ستون را از داده‌های مربوط به همان ستون کم می‌کنیم. نتایج در ماتریس (۴.۳) آورده شده‌اند.

^۴ Column effect

^۵ Change column effect

$$b_1 \begin{pmatrix} c_1 \begin{bmatrix} -1 \\ -3 \\ 2 \\ 8 \\ 5 \end{bmatrix} & a_2 \begin{bmatrix} -7 \\ -5 \\ 2 \\ 0 \\ 1 \end{bmatrix} & a_3 \begin{bmatrix} 2 \\ 5 \\ -1 \\ -4 \\ -3 \end{bmatrix} \\ c_2 \begin{bmatrix} -3 \\ -8 \\ 1 \\ 9 \\ -2 \end{bmatrix} & a_2 \begin{bmatrix} 4 \\ 0 \\ -9 \\ 7 \\ -7 \end{bmatrix} & a_3 \begin{bmatrix} -6 \\ -8 \\ 1 \\ 9 \\ 5 \end{bmatrix} \\ c_3 & & \\ c_4 & & \\ c_5 & & \end{pmatrix} \quad (4.3)$$

پس از حذف اثرات سطر و ستون، اثرات لایه را حذف می‌کنیم. در ماتریس (۴.۳) با ثابت در نظر گرفتن سطوح لایه، میانه را روی سطوح سطر و ستون محاسبه می‌کنیم. همانند قبل مقادیر اولیه اثرات لایه‌ای^۶ (LE) را برای تمام سطوح، برابر صفر در نظر می‌گیریم. اثر متغیر لایه‌ای را با^۷ CLE نشان می‌دهیم.

$$\begin{aligned} CLE_1^{(1)} &= \text{median}(-1, -7, 2, -3, 4, -6) = -2 \\ CLE_2^{(1)} &= \text{median}(-3, -5, 5, -8, 0, -8) = -4 \\ CLE_3^{(1)} &= \text{median}(2, 2, -1, 1, -9, 1) = 1 \\ CLE_4^{(1)} &= \text{median}(8, 0, -4, 9, 7, 9) = 7/5 \\ CLE_5^{(1)} &= \text{median}(5, 1, -3, -2, -7, 5) = -0/5 \\ LE_1^{(0)} &= LE_2^{(0)} = LE_3^{(0)} = LE_4^{(0)} = LE_5^{(0)} = 0 \\ LE_1^{(1)} &= CLE_1^{(1)} - LE_1^{(0)} = -2 - 0 = -2 \\ LE_2^{(1)} &= CLE_2^{(1)} - LE_2^{(0)} = -4 - 0 = -4 \\ LE_3^{(1)} &= CLE_3^{(1)} - LE_3^{(0)} = 1 - 0 = 1 \\ LE_4^{(1)} &= CLE_4^{(1)} - LE_4^{(0)} = 7/5 - 0 = 7/5 \\ LE_5^{(1)} &= CLE_5^{(1)} - LE_5^{(0)} = -0/5 - 0 = -0/5 \end{aligned}$$

میانه اثر لایه‌ای نیز از میانه‌گیری روی اثرات در تمام سطوح محاسبه می‌شود. میانه‌ی لایه‌ای

^۶ Liyer effect

^۷ Change liyer effect

مرحله اول با نماد $LE^{(1)}$ نشان داده می‌شود و به صورت زیر به دست می‌آید:

$$LE^{(1)} = \text{median}(LE_1^{(1)}, LE_2^{(1)}, LE_3^{(1)}, LE_4^{(1)}, LE_5^{(1)}) \\ = \text{median}(-2, -4, 1, 7/5, -0/5) = -0/5$$

حال با کم کردن اثر لایه‌ای از داده‌ها ماتریس (۵.۳) به دست می‌آید. این ماتریس مقدار باقی‌مانده‌های مرحله اول را ارائه می‌دهد.

$$b_1 = \begin{pmatrix} a_1 & a_2 & a_3 \\ c_1 & \begin{bmatrix} 1 \\ -5 \end{bmatrix} & \begin{bmatrix} 4 \\ -9 \end{bmatrix} \\ c_2 & \begin{bmatrix} 1 \\ -1 \end{bmatrix} & \begin{bmatrix} 9 \\ -2 \end{bmatrix} \\ c_3 & \begin{bmatrix} 1 \\ 1 \end{bmatrix} & \begin{bmatrix} -2 \\ -11/5 \end{bmatrix} \\ c_4 & \begin{bmatrix} -0/5 \\ -7/5 \end{bmatrix} & \begin{bmatrix} -11/5 \\ -2/5 \end{bmatrix} \\ c_5 & \begin{bmatrix} 5/5 \\ 1/5 \end{bmatrix} & \begin{bmatrix} -2/5 \\ 5/5 \end{bmatrix} \end{pmatrix} \quad (5.3)$$

$$b_2 = \begin{pmatrix} c_1 & \begin{bmatrix} -1 \\ 6 \end{bmatrix} & \begin{bmatrix} -4 \\ -4 \end{bmatrix} \\ c_2 & \begin{bmatrix} -4 \\ 4 \end{bmatrix} & \begin{bmatrix} -4 \\ 0 \end{bmatrix} \\ c_3 & \begin{bmatrix} 0 \\ -10 \end{bmatrix} & \begin{bmatrix} 0 \\ 1/5 \end{bmatrix} \\ c_4 & \begin{bmatrix} 1/5 \\ -0/5 \end{bmatrix} & \begin{bmatrix} 1/5 \\ 5/5 \end{bmatrix} \\ c_5 & \begin{bmatrix} -1/5 \\ -6/5 \end{bmatrix} & \begin{bmatrix} 5/5 \\ 5/5 \end{bmatrix} \end{pmatrix}$$

در نهایت مقدار میانه کل مرحله اول $\mu^{(1)}$ از مجموع سه میانه سطر، ستون و لایه حاصل می‌شود.

$$\mu^{(1)} = RE^{(1)} + RE^{(1)} + LE^{(1)} = 1/5 + 1 - 0/5 = 2$$

اگر اثرات عوامل همچنان در باقی‌مانده‌های مرحله اول مشهود بود، می‌توان میانه‌پالایش‌های مراحل دیگر را نیز ادامه داد. با این تفاوت که باقی‌مانده‌های هر مرحله، داده‌های ورودی مرحله بعد محسوب می‌شوند. تکرار مراحل تا زمانی ادامه پیدا می‌کند که محقق به هدف خود برای پالایش داده‌ها برسد و اثر عوامل به طور کامل از داده‌ها خارج شده باشد.

۲.۳ کریگینگ سه بعدی میانه پالایش

کرسی (۱۹۹۳) کریگینگ میانه‌پالایش^۸ (MPK) را ارائه داد. او از میانه پالایش دو طرفه (توکی، ۱۹۷۷؛ امرسون و هوگلین، ۲۰۰۰) برای استخراج اثرات سطح متوسطی از داده‌های فضایی و

^۸Median polish kriging

استفاده از کریگینگ معمولی (OK) برای پیشگویی فضایی فرآیند باقی‌مانده‌ها استفاده کرد. MPK با جزئیات بیشتر و کاربرد آن بر روی داده‌های زمین‌آماري که به صورت منظم یا منظم توزیع شده‌اند و همینطور بر روی داده‌های شمارشی در **کرسی (۱۹۹۳)** توضیح داده شده است. این روش نسبت به کریگینگ عام (UK) دارای دو مزیت است. در وهله اول مؤلفه‌ی میانگین در (۱۴.۱) در روش میانه پالایش نسبت به داده‌های دورافتاده مقاوم هستند و باقی‌مانده‌هایی با کمترین اریبی را برای تابع ساختار پیشگویی می‌کند (**کرسی و گلونک، ۱۹۸۴**). دوم اینکه پیشین معلومی برای تابع ساختار، تغییرنگار، فرض نمی‌شود.

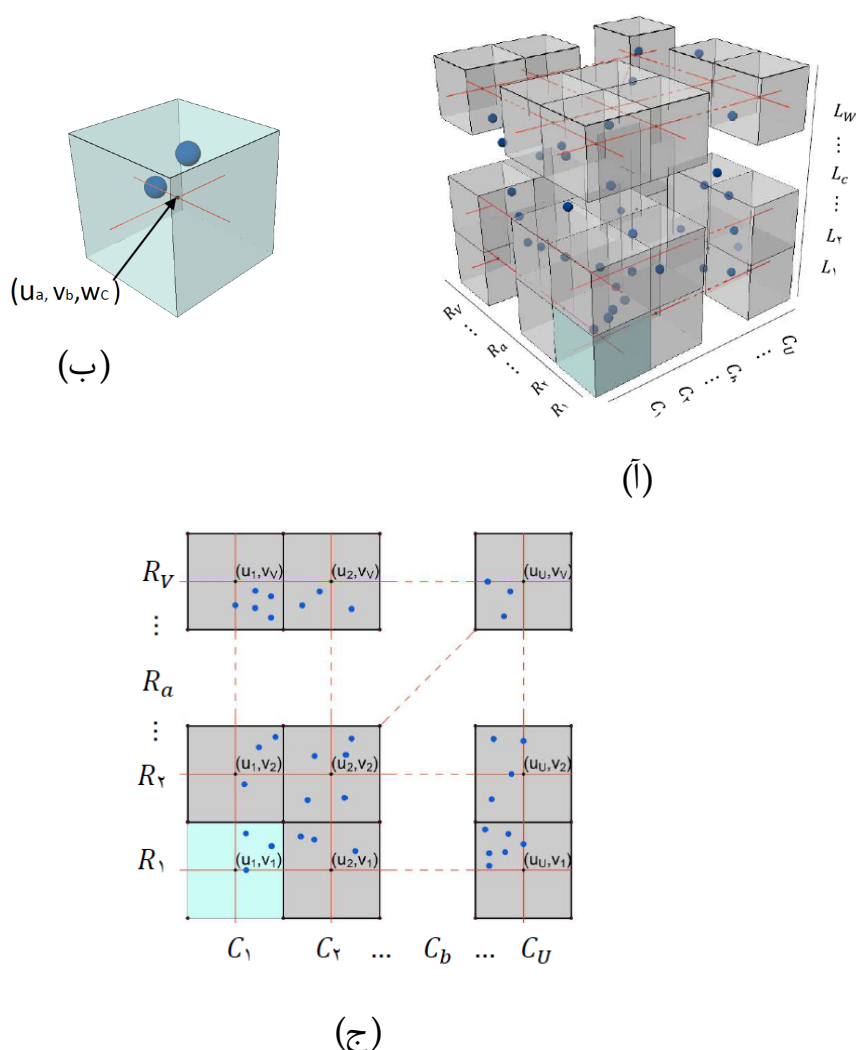
۱.۲.۳ تجزیه فرآیند فضایی

در این بخش فرض می‌کنیم، هدف تحلیل و پیشگویی فضایی در میدان $\{Z(s), s \in D\}$ است که در آن شاخص فضایی s در $D \subseteq \mathbb{R}^3$ تغییر می‌کند. این فرآیند را می‌توان به صورت زیر تجزیه کرد:

$$Z(s) = \mu(s) + \delta(s) \quad (۶.۳)$$

که در آن $\mu(\cdot)$ مؤلفه‌ی روند فضایی و $\delta(\cdot)$ مؤلفه‌ی باقی‌مانده‌های تصادفی است که نشان دهنده‌ی تغییرات تصادفی فرآیند است. ما در این پایان‌نامه از رویکرد بریک (۲۰۰۱) و مارتینز و همکاران (۲۰۱۷) برای داده‌هایی که توزیع فضایی نامنظم دارند استفاده می‌کنیم ولی تحلیل ساختار را از دو بعد به سه بعد $U \times V \times W$ گسترش می‌دهیم. فرض می‌کنیم سه بعد به صورت جداگانه در یک شکل جمعی روی متغیر پاسخ تأثیر می‌گذارند. داده‌ها را در یک آرایه سه بعدی مرتب می‌کنیم. انتظار می‌رود که مشارکت طول، عرض و لایه بر متغیر پاسخ مشخص باشد.

در اینجا R_a, C_b و L_c به ترتیب عوامل سطر، ستون و لایه، با سطح‌های U, V و W هستند. اگر مختصات فضایی داده‌ها نامنظم باشند، با توجه به فاصله‌ی بین مختصات در هر بعد، یک شبکه مکعبی $p \times q \times r$ تعریف می‌کنیم تا در منطقه‌ی تحقیق قرار گیرد. این شبکه‌ی مکعبی داده‌ها را در سلول‌های مکعبی خود جای می‌دهد. توجه کنید که بعضی از سلول‌ها ممکن است خالی باشند. سپس هر داده، مختصات نزدیک‌ترین گره را دریافت می‌کند. در داده‌های چندبعدی، اگر چند داده مختصات یک گره را اختیار کنند، باید مقادیر داده‌ها با یک مقدار مناسب مانند میانه جایگزین شود. شکل ۱.۳ شبکه‌بندی به روش میانه‌پالایش برای سه بعد و برشی از لایه، دو بعد را نشان داده است.



شکل ۱.۳: شماتیک شبکه‌بندی به روش میانه‌پالایش (آ) و (ج) سه‌بعدی، (ب) دو بعد

روند داده‌ها، $\mu(s)$ ، ممکن است به جهت مولفه‌ها و اثرات متقابل آن‌ها وابسته باشد، اما در اینجا اثرات را به صورت ساده در نظر گرفته‌ایم. پیرو روش **بریک** (۲۰۰۱) با استفاده از روش میانه پالایش در یک ساختار سه بعدی به استخراج مکرر اثرات سطر، ستون و لایه‌ها بر اساس الگوریتم بخش قبلی می‌پردازیم. شرط توقف ممکن است همگرایی باشد یا ممکن است توسط خود محقق تعریف شود. فرمول ریاضی الگوریتم میانه‌پالایش در بخش ۱.۱.۳ شرح داده شد.

برای داده‌ای در موقعیت $Z(s_{abc})$ $a = 1, \dots, U, b = 1, \dots, V, c = 1, \dots, W, s_{abc} \in D$ فرآیند تصادفی به صورت

$$Z(s_{abc}) = \mu_z(s_{abc}) + \delta_{abc} \quad (7.3)$$

است که در آن

$$\mu_z(s_{abc}) = \mu + R_a + C_b + L_c. \quad (8.3)$$

δ_{abc} تغییرات ناشی از نوسانات طبیعی و فرآیندهای اندازه‌گیری است. علاوه بر این μ میانگین کل و R_a ، C_b و L_c به ترتیب اثر سطر a ام، اثر ستون b ام و اثر لایه‌ی c ام است. سپس به حذف اثرات سه عامل می‌پردازیم تا جایی که به شرط (۱.۳) برسیم.

۲.۲.۳ درونیابی و برون‌یابی اثرات میانه پالایش

مقدار $\mu_z(s_{abc})$ داده‌شده در فرمول (۸.۳) فقط مقدار میانگین هر سلول را برآورد می‌کند. پس به نوعی از درونیابی یا برون‌یابی برای تعریف مؤلفه‌ی روند در مکان جدید پیشگویی نیاز داریم. بنابراین هر مکان فضایی $s = (u, v, w)$ را با ناحیه کراندار و ترکیبی از خطوط شبکه یعنی $u_a \leq u < u_{a+1}$ ، $v_b \leq v < v_{b+1}$ و $w_c \leq w < w_{c+1}$ در نظر می‌گیریم. به‌طوری که $a \in \{1, \dots, U\}$ ، $b \in \{1, \dots, V\}$ و $c \in \{1, \dots, W\}$. در این صورت برآورد روند فضایی به‌صورت زیر تعریف می‌شود:

$$\begin{aligned}\hat{\mu}_{mpk}(s_{abc}) = & \hat{\mu} + \hat{R}_a + \frac{u - u_a}{u_{a+1} - u_a}(\hat{R}_{a+1} - \hat{R}_a) \\ & + \hat{C}_b + \frac{v - v_b}{v_{b+1} - v_b}(\hat{C}_{b+1} - \hat{C}_b) \\ & + \hat{L}_c + \frac{w - w_c}{w_{c+1} - w_c}(\hat{L}_{c+1} - \hat{L}_c).\end{aligned}$$

در این‌جا برآورد روند میانه پالایش $\hat{\mu}_{mpk}$ حاصل از درونیابی خطی است و $\hat{L}_c$ و $\hat{C}_b$ ، $\hat{R}_a$ به ترتیب برآورد اثرهای سطری a ام، ستونی b ام و لایه‌ای c ام است. برای برآورد روند خارج از شبکه، **بریک** (۲۰۰۱) برون‌یابی را پیشنهاد کرده است. ما نیز برون‌یابی پیشنهادی او را به حالت سه بعدی گسترش می‌دهیم. حالت‌های مختلف برون‌یابی برآورد روند به‌صورت زیر هستند:

- وقتی $s = (u, v, w)$ با $u_1 < u < u_{a+1}$ ، $v_1 \leq v_b < v < v_{b+1} < v_V$ و $w_1 \leq w_c < w < w_{c+1} < w_W$ داریم

$$\begin{aligned}\hat{\mu}_{mpk}(s_{abc}) = & \hat{\mu} + \hat{R}_1 + \frac{u - u_1}{u_{a+1} - u_1}(\hat{R}_{a+1} - \hat{R}_1) \\ & + \hat{C}_b + \frac{v - v_b}{v_{b+1} - v_b}(\hat{C}_{b+1} - \hat{C}_b) \\ & + \hat{L}_c + \frac{w - w_c}{w_{c+1} - w_c}(\hat{L}_{c+1} - \hat{L}_c).\end{aligned}$$

- وقتی $s = (u, v, w)$ با $u_1 \leq u_a < u < u_{a+1} < u_U$ ، $v < v_1$ و $w_1 \leq w_c < w < w_{c+1} < w_W$ داریم

$$\begin{aligned}\hat{\mu}_{mpk}(s_{abc}) = & \hat{\mu} + \hat{R}_a + \frac{u - u_a}{u_{a+1} - u_a}(\hat{R}_{a+1} - \hat{R}_a) \\ & + \hat{C}_1 + \frac{v - v_1}{v_2 - v_1}(\hat{C}_2 - \hat{C}_1) \\ & + \hat{L}_c + \frac{w - w_c}{w_{c+1} - w_c}(\hat{L}_{c+1} - \hat{L}_c).\end{aligned}$$

- وقتی $s = (u, v, w)$ با $u_1 \leq u_a < u < u_{a+1} < u_U$ ، $v_1 \leq v_b < v < v_{b+1} < v_V$ و $w < w_1$ داریم

$$\begin{aligned}\hat{\mu}_{mpk}(s_{abc}) = & \hat{\mu} + \hat{R}_a + \frac{u - u_a}{u_{a+1} - u_a}(\hat{R}_{a+1} - \hat{R}_a) \\ & + \hat{C}_b + \frac{v - v_b}{v_{b+1} - v_b}(\hat{C}_{b+1} - \hat{C}_b) \\ & + \hat{L}_1 + \frac{w - w_1}{w_2 - w_1}(\hat{L}_2 - \hat{L}_1).\end{aligned}$$

- وقتی $s = (u, v, w)$ با $u < u_1$ ، $v < v_1$ و $w_1 \leq w_c < w < w_{c+1} < w_W$ داریم

$$\begin{aligned}\hat{\mu}_{mpk}(s_{abc}) = & \hat{\mu} + \hat{R}_1 + \frac{u - u_1}{u_2 - u_1}(\hat{R}_2 - \hat{R}_1) \\ & + \hat{C}_1 + \frac{v - v_1}{v_2 - v_1}(\hat{C}_2 - \hat{C}_1) \\ & + \hat{L}_c + \frac{w - w_c}{w_{c+1} - w_c}(\hat{L}_{c+1} - \hat{L}_c).\end{aligned}$$

- وقتی $s = (u, v, w)$ با $u < u_1$ ، $v_1 \leq v_b < v < v_{b+1} < v_V$ و $w < w_1$ داریم

$$\begin{aligned}\hat{\mu}_{mpk}(s_{abc}) = & \hat{\mu} + \hat{R}_1 + \frac{u - u_1}{u_2 - u_1}(\hat{R}_2 - \hat{R}_1) \\ & + \hat{C}_b + \frac{v - v_b}{v_{b+1} - v_b}(\hat{C}_{b+1} - \hat{C}_b) \\ & + \hat{L}_1 + \frac{w - w_1}{w_2 - w_1}(\hat{L}_2 - \hat{L}_1).\end{aligned}$$

- وقتی $s = (u, v, w)$ با $u_1 \leq u_a < u < u_{a+1} < u_U$ ، $v < v_1$ و $w < w_1$ داریم

$$\begin{aligned}\hat{\mu}_{mpk}(s_{abc}) = & \hat{\mu} + \hat{R}_a + \frac{u - u_a}{u_{a+1} - u_a}(\hat{R}_{a+1} - \hat{R}_a) \\ & + \hat{C}_1 + \frac{v - v_1}{v_2 - v_1}(\hat{C}_2 - \hat{C}_1) \\ & + \hat{L}_1 + \frac{w - w_1}{w_2 - w_1}(\hat{L}_2 - \hat{L}_1).\end{aligned}$$

- وقتی $s = (u, v, w)$ با $u < u_1$ ، $v < v_1$ و $w < w_1$ داریم

$$\begin{aligned}\hat{\mu}_{mpk}(s_{abc}) = & \hat{\mu} + \hat{R}_1 + \frac{u - u_1}{u_2 - u_1}(\hat{R}_2 - \hat{R}_1) \\ & + \hat{C}_1 + \frac{v - v_1}{v_2 - v_1}(\hat{C}_2 - \hat{C}_1) \\ & + \hat{L}_1 + \frac{w - w_1}{w_2 - w_1}(\hat{L}_2 - \hat{L}_1).\end{aligned}$$

- وقتی $s = (u, v, w)$ با $u_U < u$ ، $v_1 \leq v_b < v < v_{b+1} < v_V$ و $w_1 \leq w_c < w < w_{c+1} < w_W$ داریم

w_W ، داریم

$$\begin{aligned}\hat{\mu}_{mpk}(s_{abc}) = & \hat{\mu} + \hat{R}_{U-1} + \frac{u - u_{U-1}}{u_U - u_{U-1}}(\hat{R}_U - \hat{R}_{U-1}) \\ & + \hat{C}_b + \frac{v - v_b}{v_{b+1} - v_b}(\hat{C}_{b+1} - \hat{C}_b) \\ & + \hat{L}_c + \frac{w - w_c}{w_{c+1} - w_c}(\hat{L}_{c+1} - \hat{L}_c).\end{aligned}$$

به‌طور مشابه برای حالت‌های $v_V < v$ و $w_W < w$ عمل می‌کنیم.

۳.۲.۳ کریگینگ معمولی روی باقی‌مانده‌ها

در نهایت پیشگویی توسط کریگینگ میانه پالایش از مجموع برآورد روند میانه پالایش با مقدار پیشگویی حاصل از کریگینگ معمولی روی باقی‌مانده‌های به‌دست آمده از میانه پالایش در همان سطح به‌دست می‌آید. یعنی مقدار پیشگویی میدان در نقطه مکانی جدید s_0 عبارت است از

$$\hat{Z}_{mpk}(s_0) = \hat{\mu}_{mpk}(s_0) + \hat{\delta}_{mpk}(s_0)$$

به‌طوری که

$$\hat{\delta}_{mpk}(s_0) = \hat{\theta}_\delta + \mathbf{C}_\delta^T \Sigma_\delta^{-1} (\delta - \mathbf{1} \hat{\theta}_\delta) \quad s_0 \in \mathbb{R}^3$$

که در آن $\mathbf{Z}_{mpk} = (Z_{mpk}(s_1), \dots, Z_{mpk}(s_n))^T$ با $\delta = \mathbf{Z}_{mpk} - \mu_{mpk}$ ، $\Sigma_\delta = \text{Cov}(\delta)$ ، $\mathbf{1} = (1, \dots, 1)^T$ و $\delta = (\delta(s_1), \dots, \delta(s_n))^T$ ، $\delta_0 = \delta(s_0)$ با $\mathbf{C}_\delta = \text{Cov}(\delta, \delta_0)$ ، $\hat{\mu}_{mpk} = (\hat{\mu}_{mpk}(s_1), \dots, \hat{\mu}_{mpk}(s_n))^T$ و $\hat{\theta}_\delta = (\mathbf{1}^T \Sigma_\delta^{-1} \mathbf{1})^{-1} \mathbf{1}^T \Sigma_\delta^{-1} \delta$ ثابت، برای فرآیند باقی‌مانده‌های $\delta(\cdot)$ و بر اساس ساختار فضایی ماتریس کوواریانس Σ_δ است. کمترین توان‌های دوم خطای پیشگو^۹ (MSPE) برای کریگینگ میانه‌پالایش نیز برابر با مقدار واریانس کریگینگ معمولی باقی‌مانده‌ها است. بنابراین

$$\begin{aligned}\text{MSPE}(Z(s_0), \hat{Z}_{mpk}(s_0)) & \approx \text{MSPE}(\delta(s_0), \hat{\delta}_{mpk}(s_0)) \\ & = \hat{\sigma}_{mpk}^2(s_0) = \sigma^2 - \mathbf{C}_\delta^T \Sigma_\delta^{-1} \mathbf{C}_\delta + (\mathbf{1} - \mathbf{1}^T \Sigma_\delta^{-1} \mathbf{C}_\delta)^2 / (\mathbf{1}^T \Sigma_\delta^{-1} \mathbf{1})\end{aligned}$$

که در آن $\sigma^2 = \text{Cov}(0)$ همان ازاره مدل برازش داده‌شده به تغییرنگار است.

به‌منظور ارزیابی عملکرد رهیافت پیشنهادی ما بر اساس روش میانه‌پالایش، نتایج حاصل از پیشگویی را در قالب یک مطالعه شبیه‌سازی با روش کریگینگ عام که در آن برآورد تابع روند با روش اسپلین محاسبه می‌شود، مقایسه کردیم. برای این منظور ابتدا روش برآورد ناپارامتری یک تابع روند با روش اسپلین را معرفی می‌کنیم.

^۹Mean squared prediction error

۳.۳ روش اسپالین

در تحلیل رگرسیونی داده‌های فضایی، روش مرسوم استفاده از مدل‌هایی با اثرات آمیخته خطی پارامتری است، به‌طوری‌که وابستگی فضایی به‌عنوان یک اثر تصادفی فضایی در نظر گرفته می‌شود و معمولاً یک فرآیند گوسی با میانگین صفر و تابع کوواریانس پارامتری فرض می‌شود (استروپ، ۲۰۱۲). بسیاری از نویسندگان از جمله ریپلی (۱۹۸۱) و کرسی (۱۹۹۳) از روش‌های پارامتری برای استنباط آماری در مدل‌های فضایی استفاده کرده‌اند. با این حال اغلب اوقات، با شرایطی رو به رو می‌شویم که یک مدل پارامتری خاص را با اطمینان نمی‌توان به‌کار گرفت. بنابراین باید از یک روش ناپارامتری جایگزین استفاده کرد.

روش اسپالین می‌تواند به‌عنوان یک روش جایگزین برای مدل‌بندی تابع روند فضایی مورد استفاده قرار گیرد. این روش، یک روش نیمه‌پارامتری به‌وجود می‌آورد که برای محاسبه‌ی اثرات روند فضایی نیازی به مشخصات ساختاری پارامتری ندارد (لودویگ و همکاران، ۲۰۲۰). یکی از رایج‌ترین روش‌های هموارسازی رگرسیون ناپارامتری، اسپالین‌ها هستند که برای اولین بار توسط شوئنبرگ (۱۹۴۶) معرفی شد. وی برای نخستین بار قضیه‌ای برای اثبات وجود اسپالین درون‌یاب مطرح کرد و بعد مطالعه‌ی چندجمله‌ای‌های تکه‌ای برای توابع اسپالین آغاز شد. با فرض $s \in \mathbb{R}$ ، فاصله‌ی $[a, b]$ را می‌توان به‌صورت اجتماع دو مجموعه $[a, t] \cup [t, b]$ در نظر گرفت. اسپالین مرتبه‌ی یک، تابعی است که در هر یک از بازه‌های فوق یک تابع خطی و در نقطه‌ی t پیوسته باشد. اسپالین خطی مرتبه اول را می‌توان به‌صورت

$$\mu(s) = \begin{cases} \beta_0 + \beta_1 s & s \leq t \\ \beta'_0 + \beta'_1 s & s > t \end{cases} \quad (9.3)$$

نمایش داد. با در نظر گرفتن شرط پیوستگی تابع μ در نقطه‌ی $s = t$ داریم

$$\begin{aligned} \beta_0 + \beta_1 t &= \beta'_0 + \beta'_1 t \\ \beta'_0 &= \beta_0 + \beta_1 t - \beta'_1 t \end{aligned} \quad (10.3)$$

و با جایگذاری β'_0 در ضابطه‌ی دوم μ داریم

$$\beta'_0 + \beta'_1 s = \beta_0 + \beta_1 t - \beta'_1 t + \beta'_1 s = \beta_0 + \beta_1 t + \beta'_1 (s - t)$$

پس

$$\mu(s) = \begin{cases} \beta_0 + \beta_1 s & s \leq t \\ \beta_0 + \beta_1 t + \beta'_1 (s - t) & s > t \end{cases} \quad (11.3)$$

که می‌توان تابع دو ضابطه‌ای را به‌صورت

$$\mu(s) = \beta_0 + \beta_1 s + \beta_2 (s - t)_+$$

نوشت که در آن $(s - t)_+ = \mu(s)(s - t, 0) = (s - t)I(s > t)$ ، تابع نشانگر و $\beta_2 = \beta'_1 - \beta_1$ اختلاف بین شیب خطوط اول و دوم تابع را نشان می‌دهد. به عبارتی می‌توان به این نتیجه رسید که تابع μ ترکیب خطی از توابع پایه ۱، s و $(s - t)_+$ است. اسپلاین مرتبه اول را اسپلاین خطی نیز می‌نامند.

اگر $[a, b]$ به $k + 1$ زیربازه افراز شده باشد، یعنی

$$[a, t_1), [t_1, t_2), \dots, [t_K, b]$$

به‌طوری که $a < t_1 < \dots < t_K < b$ ، نقاط $\{t_i\}_{i=1}^K$ را گره می‌نامند. با این شرایط اسپلاین مرتبه اول را می‌توان به‌صورت

$$\mu(s) = \beta_0 + \beta_1 s + \sum_{k=1}^K \beta_{k+1} (s - t_k)_+$$

نشان داد که ترکیب خطی از توابع پایه ۱، s ، $(s - t_1)_+$ ، $\dots$ و $(s - t_K)_+$ است. به‌طور مشابه، توابع اسپلاین مرتبه‌های بالاتر تعریف می‌شوند. اگر هر یک از بازه‌های مذکور یک چندجمله‌ای از مرتبه d باشد که $\mu(\cdot)$ و مشتق‌های آن از مرتبه‌ی ۱ تا $d - 1$ در گره‌ها پیوسته باشند، آن‌گاه تابع $\mu(\cdot)$ را یک اسپلاین مرتبه d می‌نامند. تابع اسپلاین مرتبه d را می‌توان به‌صورت

$$\mu(s) = \beta_0 + \beta_1 s + \beta_2 s^2 + \dots + \beta_d s^d + \sum_{k=1}^K \beta_{k+d} (s - t_k)_+^d$$

بیان کرد، که در آن $(s - t_k)_+^d = (s - t_k)^d I(s > t_k)$ تابع بریده‌شده توانی از مرتبه‌ی d است. به بیان ساده‌تر، این تابع ترکیب خطی از یک چندجمله‌ای مرتبه d و توابع بریده‌شده است. مجموعه‌ی

$$\{1, s, s^2, \dots, s^d, (s - t_1)_+^d, \dots, (s - t_K)_+^d\}$$

یک پایه برای فضای اسپلاین مرتبه d با گره‌های $t_1, t_2, \dots, t_K$ است. اگر $d = 3$ باشد، تابع مورد نظر را اسپلاین مکعبی^{۱۰} می‌نامند. اسپلاین مکعبی به دلیل انعطاف‌پذیری بالایی که دارد یکی از رایج‌ترین و پرکاربردترین نوع اسپلاین‌هاست که در بسیاری از علوم مورد استفاده قرار می‌گیرد.

مدل رگرسیون ناپارامتری $Z = \mu(s) + \delta$ را در نظر بگیرید، که در آن

$$\mathbf{Z} = (Z_1, \dots, Z_n)^T$$

$$\boldsymbol{\mu}(s) = (\mu(s_1), \dots, \mu(s_n))^T \quad (12.3)$$

$$\boldsymbol{\delta} = (\delta_1, \dots, \delta_n)$$

^{۱۰}Cubic spline

اگر به جای تابع $\mu(\cdot)$ از تابع اسپالین استفاده شود، رگرسیون به دست آمده را رگرسیون اسپالین می نامند. بنابراین می توان مدل زیر را ارائه کرد:

$$Z_i = \beta_0 + \beta_1 s_i + \beta_2 s_i^2 + \dots + \beta_d s_i^d + \sum_{k=1}^K \beta_{d+k} (s_i - t_k)_+^d + \delta_i \quad i = 1, \dots, n$$

که در آن δ_i در پذیره های زیربنایی رگرسیون صدق می کند. با استفاده از نمادگذاری ماتریسی، رگرسیون اسپالین به صورت

$$\mathbf{Z} = \mathbf{B}\beta + \delta$$

نمایش داده می شود که در آن $\beta = (\beta_0, \dots, \beta_{d+K})$ و ماتریس طرح $\mathbf{B}$ به صورت

$$\mathbf{B} = \begin{pmatrix} 1 & s_1 & s_1^2 & \dots & s_1^d & (s_1 - t_1)_+^d & (s_1 - t_2)_+^d & \dots & (s_1 - t_K)_+^d \\ 1 & s_2 & s_2^2 & \dots & s_2^d & (s_2 - t_1)_+^d & (s_2 - t_2)_+^d & \dots & (s_2 - t_K)_+^d \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 1 & s_n & s_n^2 & \dots & s_n^d & (s_n - t_1)_+^d & (s_n - t_2)_+^d & \dots & (s_n - t_K)_+^d \end{pmatrix}$$

$n \times (d+K+1)$
(۱۳.۳)

است. در $\mathbf{B}$ درایه های ستون اول همگی ۱ هستند و درایه های d ستون بعدی مربوط به مشاهدات متغیرهای s ، s^2 ، ... و s^d و درایه های ستون های $d+2$ تا $d+K+1$ مربوط به مشاهدات حاصل از متغیرهای $(s - t_1)_+^d$ ، ... و $(s - t_K)_+^d$ به ازای گره های از پیش تعیین شده t_1 ، t_2 ، ... و t_K است. رگرسیون اسپالین تعریف شده متعلق به رده رگرسیون خطی است. به عبارتی تابع $\mu(\cdot)$ نسبت به پارامترها خطی است و هزینه های محاسباتی برآورد و منحنی رگرسیون اندک است. پس می توان نتیجه گرفت که برازش یک مدل رگرسیون اسپالین بر روی داده ها را به راحتی می توان با استفاده از روش های مرسوم پیاده سازی کرد.

با استفاده از روش کمترین توان های دوم خطا، برآوردگر بردار β به صورت $\hat{\beta} = (\mathbf{B}^T \mathbf{B})^{-1} \mathbf{B}^T \mathbf{Z}$ به دست می آید. عیبی که به این برآوردگر وارد است، هم خطی بین پارامترهای توان بریده شده

$$(s - t_1)_+^d, (s - t_2)_+^d, \dots, (s - t_K)_+^d$$

است که موجب عدم وارون پذیری یا بدشرطی ماتریس $\mathbf{B}^T \mathbf{B}$ می شود و این موضوع ناپایداری ضرایب رگرسیونی را در پی دارد.

یک روش مناسب پیشنهاد شده برای حل این مشکل، استفاده از B-اسپالین به جای پایه های توان بریده شده $(s - t_i)_+^d$ ، $i = 1, \dots, K$ است. بازه ی $[a, b]$ و تعداد K گره به صورت

$$a = t_0 < t_1 < \dots < t_K < t_{K+1} = b$$

را در نظر بگیرید. متداول ترین روش محاسبه ی پایه های B-اسپالین مرتبه ی d استفاده از پایه های B-اسپالین مرتبه ی $d-1$ است که بر اساس رابطه ی بازگشتی زیر است:

$$B_{i,d}(s) = \left(\frac{s - t_i}{t_{i+d} - t_i}\right) B_{i,d-1}(s) + \left(\frac{t_{i+d+1} - s}{t_{i+d+1} - t_{i+1}}\right) B_{i+1,d-1}(s), \quad i = -d, -d+1, \dots, K$$

که در آن $a = T_0 = t_{-1} = t_{-2} = \dots = t_{-d} = t_{-d+1}$ و $t_{K+1} = b$. تعداد پایه‌های B-اسپلاین مرتبه‌ی d برابر $d + K + 1$ است. با استفاده از این پایه‌ها می‌توان چندجمله‌ای‌های متعامد را به‌دست آورد.

در این پایان‌نامه از روش رگرسیون اسپلاین و بر اساس توابع B-اسپلاین در کنار رهیافت میانه‌پالایش، تابع روند را برآورد و کریگینگ مشابه کریگینگ میانه‌پالایش اجرا می‌کنیم. نتایج این دو روش را با هم مقایسه می‌کنیم.

۴.۳ مطالعه شبیه‌سازی

در این بخش عملکرد کریگینگ میانه‌پالایش را با داده‌های شبیه‌سازی شده سه بعدی منظم بررسی می‌کنیم. سپس این روش را با کریگینگ عامی که روند آن با استفاده از روش اسپلاین مدل شده است، مقایسه می‌کنیم.

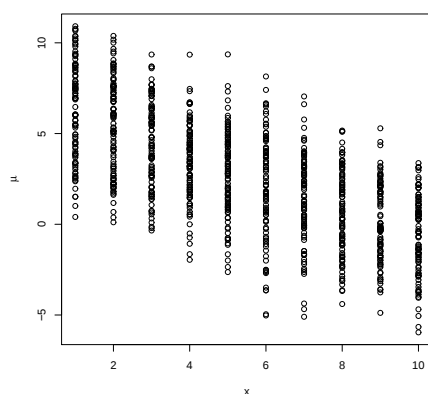
در مطالعه شبیه‌سازی دو حالت خطی و غیرخطی را برای روند واقعی داده‌ها در نظر گرفتیم که به تفکیک آن‌ها را تشریح می‌کنیم.

۱.۴.۳ روند خطی

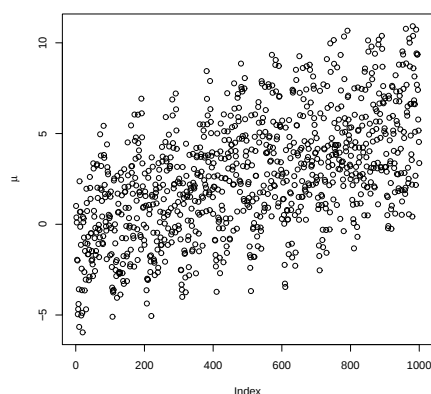
برای تولید ۱۰۰۰ داده‌ی فضایی با مختصات سه بعدی (طول، عرض و لایه) از بسته geoR در محیط نرم‌افزار R استفاده کردیم. در مدل (۶.۳) داده‌ها را از رابطه‌ی زیر بر اساس یک تابع روند خطی تولید کردیم:

$$\mu = 1 - (0.8 \times x) + (0.5 \times y) + (0.6 \times z)$$

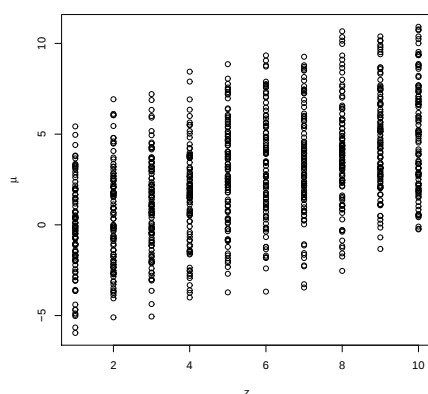
اثر تصادفی فضایی $\delta(s)$ را از یک میدان تصادفی گاوسی با میانگین صفر و ماتریس کواریانس که از مدل نمایی (۷.۱) به دست می‌آید تولید کردیم، که در آن پارامتر مدل $\phi = 0.5$ انتخاب شد. اثر قطعه‌ای، دامنه و ابزار را نیز به ترتیب برابر با ۰، ۱ و ۱ در نظر گرفتیم. نمودارهای پراکنش داده‌های خام و داده‌ها بر حسب سه محور مختصات در شکل ۲.۳ نمایش داده شده‌اند.



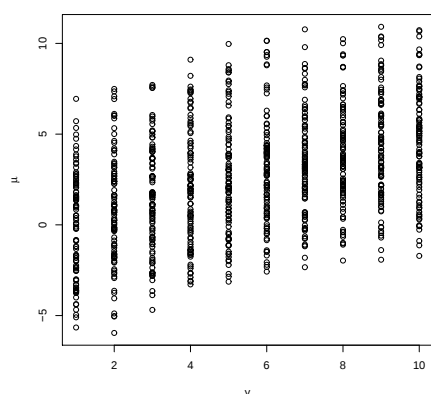
(ب)



(ا)



(د)



(ج)

شکل ۲.۳: نمودار پراکنش داده‌های خام و پراکنش نسبت به مختصات x ، y و z به ترتیب (ا)، (ب)، (ج) و (د).

نمودار نیم‌تغییرنگار تجربی و مدل برازش داده شده بر اساس سه رهیافت مختلف در شکل ۳.۳ آورده شده است. تغییرنگار برآورد شده بر اساس داده‌های اصلی (مدل خطی) در شکل ۲.۳ (ا) نشان می‌دهد که میدان مانا نیست و استفاده از داده‌ها قبل از حذف روند بیانگر خطای زیادی در برآورد ساختار وابستگی فضایی داده‌ها است. بنابراین وجود روند در داده‌ها کاملاً تایید می‌شود.

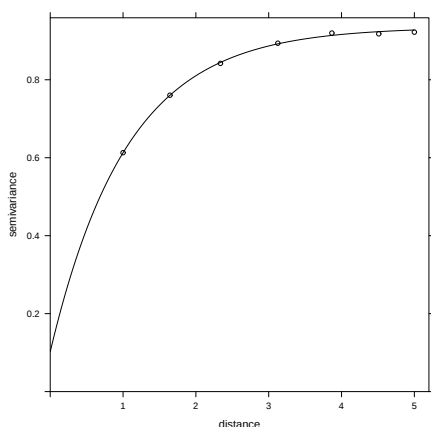
پس داده‌ها باید ابتدا روندزدایی شوند. قبل از این کار برای ارزیابی عملکرد روش‌های پیشنهادی داده‌ها را به سه روش ۵ بسته^{۱۱}، ۱۰ بسته^{۱۲} و LOOCV^{۱۳} (فریدمن و همکاران، ۲۰۰۱) به دو مجموعه آموزشی و آزمون تقسیم کردیم و با داده‌های آموزشی دو الگوریتم میانه‌پالایش و

^{۱۱} 5 Fold

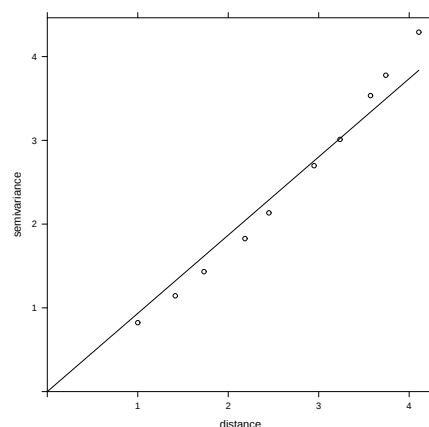
^{۱۲} 10 Fold

^{۱۳} Leave one out cross validation

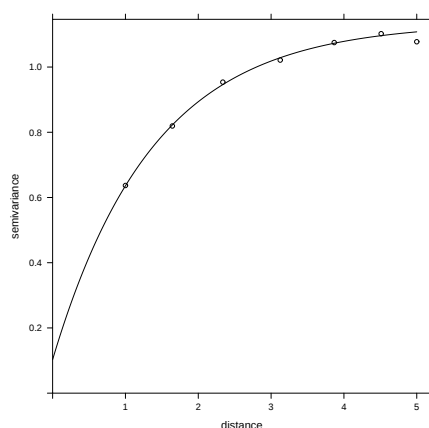
رگرسیون خطی را روی سه مختصات x ، y و z برای برآورد روند و محاسبه باقی‌مانده‌ها اجرا کردیم. نتیجه برآورد تغییرنگار بر اساس باقی‌مانده‌ها برای دو روش، مدل مشابه نمایی را پیشنهاد می‌دهد.



(ب)



(ا)

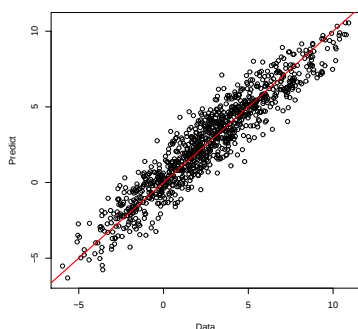


(ج)

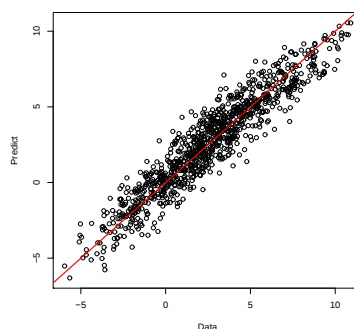
شکل ۳.۳: نمودار سه بعدی نیم‌تغییرنگار و (ا): منحنی مدل پارامتری خطی برازش داده شده روی داده‌های اصلی. (ب) منحنی مدل پارامتری نمایی برازش داده شده روی باقی‌مانده‌های حاصل از میانه‌پالایش. (ج) منحنی مدل پارامتری نمایی برازش داده شده روی باقی‌مانده‌های مدل برازش داده شده با استفاده از رگرسیون خطی

محور افقی نمودارهای پراکنش در شکل ۴.۳ نشان‌دهنده‌ی مقادیر واقعی داده‌های آزمون و محور عمودی آن‌ها بیانگر مقادیر پیشگویی به‌دست آمده از کریگینگ میانه‌پالایش و عام است. نمودارهای سمت راست نتایج حاصل از روش میانه‌پالایش و سمت چپ روش کریگینگ عام را نشان می‌دهند. در این حالت (روند خطی) در هر دو روش توزیع پراکنش داده‌ها حول نیم‌ساز ربع اول و سوم است و بیانگر مناسب بودن مدل و مقدار ناچیز باقی‌مانده‌ها است. همچنین نیم‌ساز بین داده‌ها قرار دارد و می‌توان نتیجه گرفت میانگین خطا برابر با صفر است.

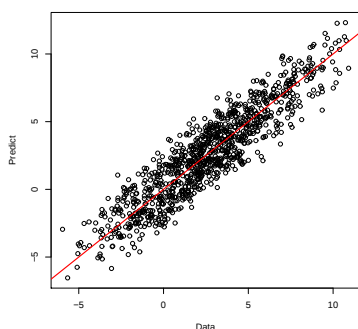
با استفاده از روش‌های اعتبارسنجی متقابل ۵ بسته، ۱۰ بسته و LOOCV مقادیر MSPE را برای دو روش محاسبه کردیم که در جدول ۱.۳ گزارش شده است.



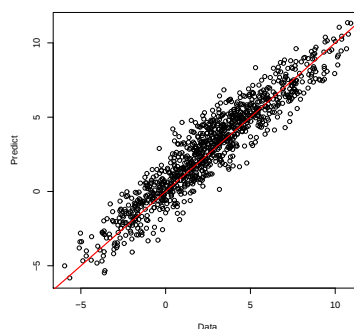
(ا)



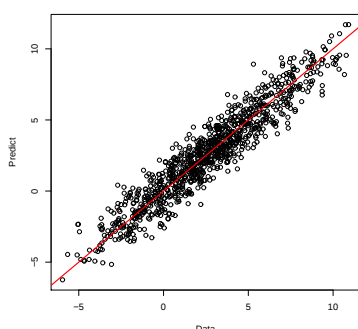
(ب)



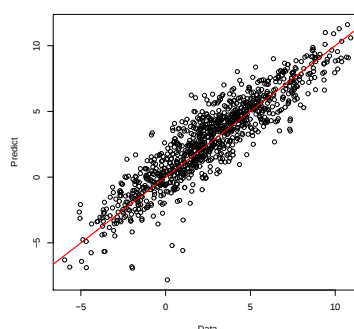
(ج)



(د)



(ه)



(و)

شکل ۴.۳: مقادیر پیشگویی در مقابل مقادیر واقعی داده‌های آزمون با استفاده از کریگینگ میانه‌پالایش (ا) ۵ بسته، (ب) ۱۰ بسته و LOOCV و با استفاده از کریگینگ عام (ج) ۵ بسته، (د) ۱۰ بسته و (و) LOOCV.

جدول ۱.۳: مقادیر MSPE

کریگینگ میانه‌پالایش			کریگینگ عام		
LOOCV	۵ بسته	۱۰ بسته	LOOCV	۵ بسته	۱۰ بسته
۱.۶	۱.۱۸	۱.۲۰	۱.۱۱	۱.۹۱	۱.۹۶

بر اساس مقادیر جدول، مشخص است که روش میانه‌پالایش با استفاده از روش‌های ۵ بسته و ۱۰ بسته برتر از روش کریگینگ عام است، اما با استفاده از اعتبارسنجی متقابل LOOCV روش کریگینگ عام بهتر عمل می‌کند. نمودارهای پراکنش شکل ۴.۳ در دو رویکرد اعتبارسنجی متقابل نیز این موضوع را تأیید می‌کنند.

۲.۴.۳ روند غیر خطی

در شبیه‌سازی دوم روند داده‌ها را غیرخطی در نظر گرفتیم و از رابطه‌ی زیر آن‌ها را تولید کردیم:

$$\mu = 1 + 0.8 \log(x) + 0.6 \sin(y) - 0.8 \cos(z) + 0.5z \times (-0.5)y \quad (14.3)$$

اثر تصادفی فضایی را مشابه روند خطی تولید کردیم. نمودارهای پراکنش داده‌های خام و به تفکیک سه بعد در شکل ۵.۳ آورده شده‌اند.

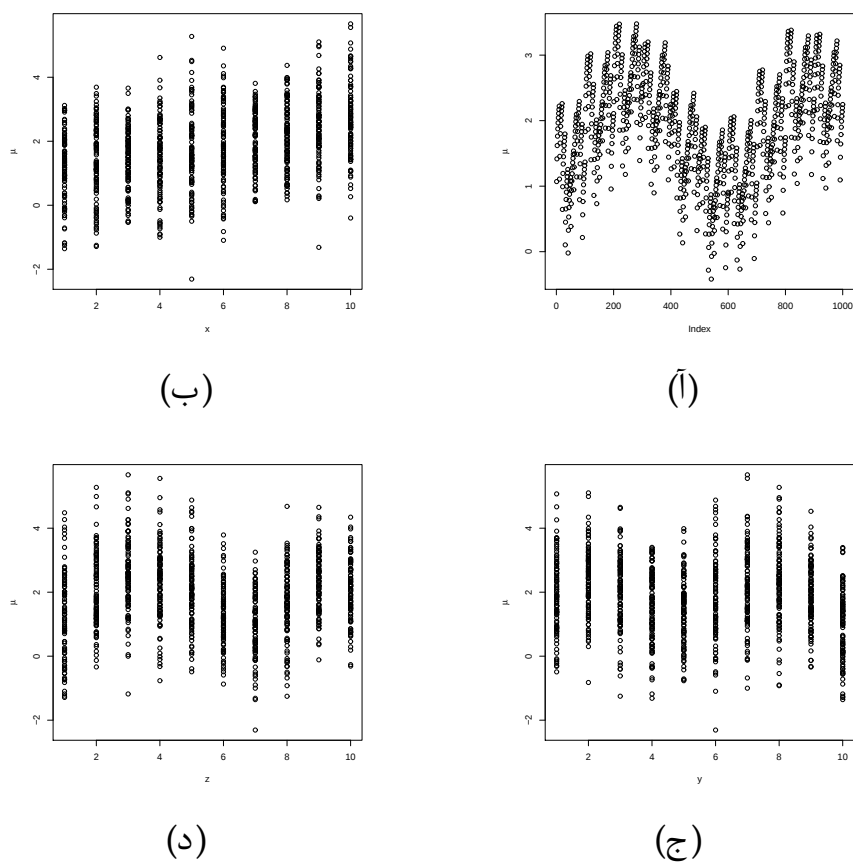
با اجرای کریگینگ سه بعدی عادی روی باقی‌مانده‌های حاصل از سه روش رگرسیون خطی، میانه‌پالایش و اسپلین، نتایج اعتبارسنجی متقابل پیشگویی داده‌های آزمون را محاسبه و نمودارهای پراکنش آن‌ها در شکل ۶.۳ گزارش کردیم.

مشابه شبیه‌سازی اول، ابتدا پس از حذف روند تابع تغییرنگار بر اساس هر سه روش میانه‌پالایش کریگینگ عام با روند خطی و کریگینگ عام با روند ناپارامتری حاصل از اسپلین برآورد شدند و کریگینگ معمولی روی باقی‌مانده‌های حاصل پیاده‌سازی شد. نتایج پیشگویی برای داده‌های مجموعه‌های آزمون حاصل از سه روش ۵ بسته، ۱۰ بسته و LOOCV در شکل ۶.۳ گزارش شده‌اند. در این حالت نیز توزیع پراکنش داده‌ها نسبتاً حول نیم‌ساز ربع اول و سوم است و بیانگر مناسب بودن روش‌ها است.

مقدار MSPE را برای سه روش مذکور محاسبه کردیم که در جدول ۲.۳ گزارش شده‌اند.

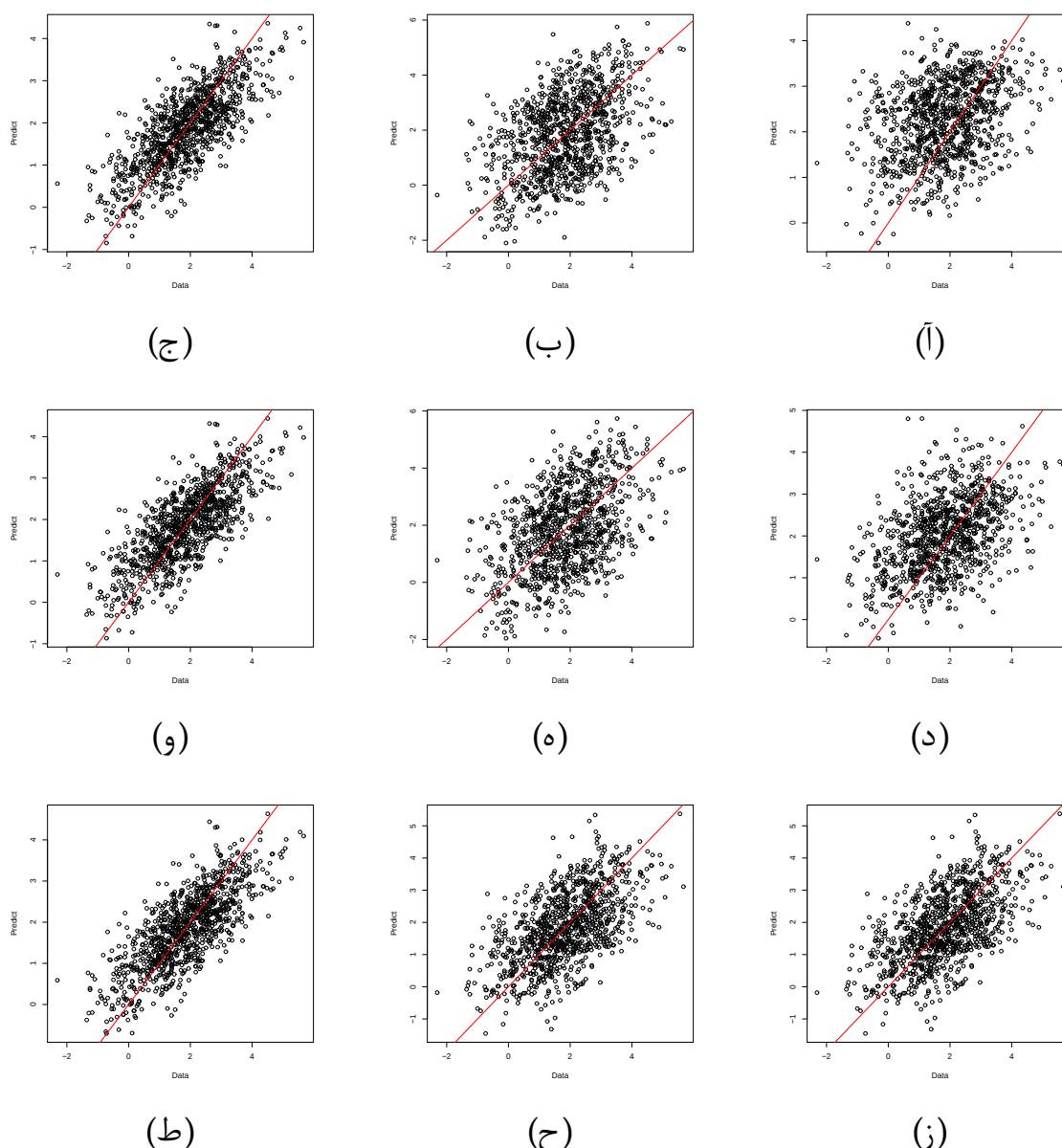
جدول ۲.۳: مقادیر MSPE

کریگینگ عام با روند اسپلین			کریگینگ میانه‌پالایش			کریگینگ عام		
LOOCV	۵ بسته	۱۰ بسته	LOOCV	۵ بسته	۱۰ بسته	LOOCV	۵ بسته	۱۰ بسته
۰.۶۱۸	۰.۶۲۴	۰.۶۳۵	۱.۲۱	۱.۳۴	۱.۶۵	۱.۱۳	۱.۹۵	۱.۹۸



شکل ۵.۳: نمودار پراکنش داده‌ها و پراکنش نسبت به مختصات x ، y و z به ترتیب (ا)، (ب)، (ج) و (د).

بر اساس مقادیر جدول، مشابه حالت روند خطی، روش میان‌پالایش در پیشگویی فضایی در دو روش ۵ بسته و ۱۰ بسته بهتر از کریگینگ عام عمل می‌کند ولی در روش LOOCV، کریگینگ عام عملکرد بهتری نسبت به کریگینگ میان‌پالایش دارد. البته با توجه به مقادیر جدول می‌توان درک کرد که روش اسپلاین بهتر از دو رقیب دیگر عمل کرده است، که نشان از توانایی بالای این روش در برآورد توابع غیرخطی دارد. نمودارهای پراکنش شکل ۶.۳ در سه رویکرد اعتبارسنجی متقابل نیز این موضوع را تأیید می‌کنند.



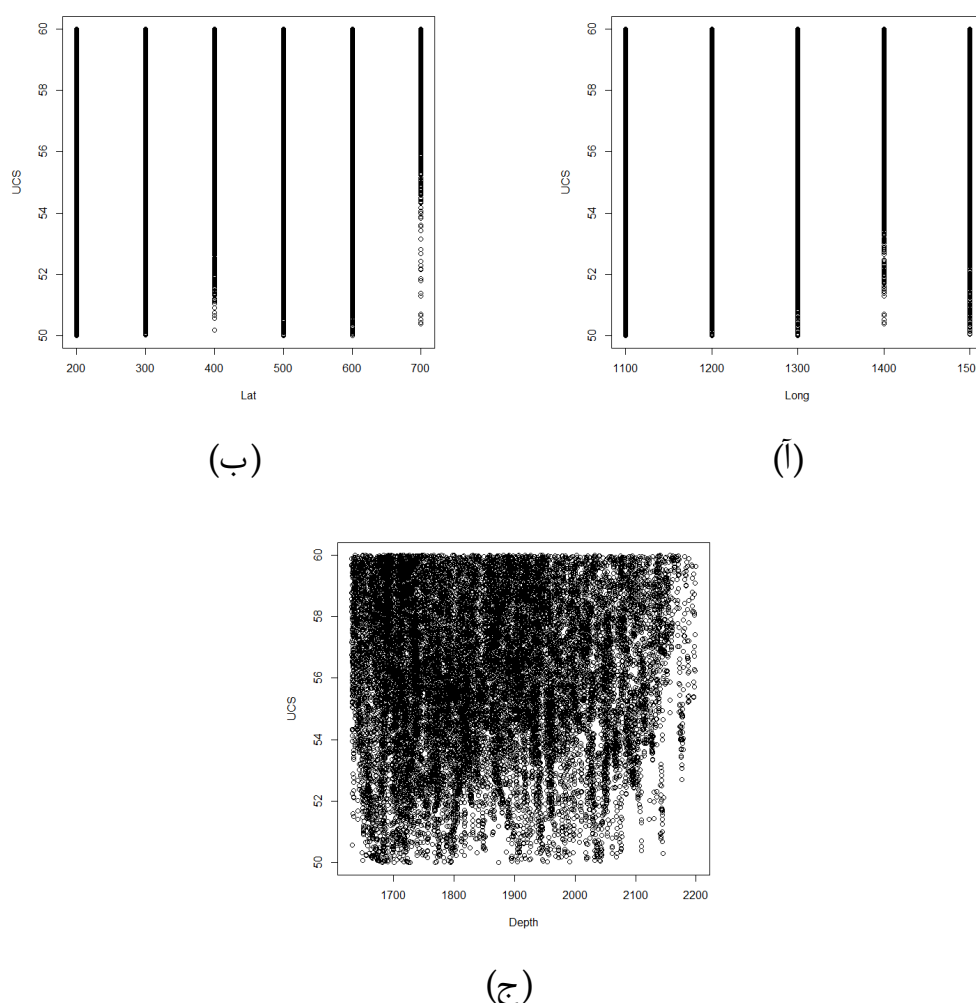
شکل ۵.۳: مقادیر پیشگویی در مقابل مقادیر واقعی داده‌های آزمون با استفاده از کریگینگ میانه‌پالایش (آ) ۵ بسته، (د) ۱۰ بسته و (ز) LOOCV، کریگینگ عام با روند خطی (ب) ۵ بسته، (ه) ۱۰ بسته و (ح) LOOCV و با استفاده از کریگینگ عام روی باقی‌مانده‌های حاصل از اسپلین (ج) ۵ بسته، (و) ۱۰ بسته و (ط) LOOCV.

۵.۳ مطالعه پیشگویی سه بعدی ناپارامتری مقاومت فشاری تک محوری سنگ

در اینجا به پیشگویی سه بعدی مقاومت فشاری تک محوری سازند با استفاده از روش‌های ناپارامتری می‌پردازیم. با به‌کارگیری از رویکرد ارزیابی متقابل، داده‌های UCS را به‌طور تصادفی

به دو مجموعه‌ی آزمایشی و آزمون تقسیم کردیم. پیش‌تر گفته شد که میدان تصادفی را می‌توان به صورت $Z(s) = \mu(s) + \delta(s)$, $s \in D$ تجزیه کرد، پس ابتدا روند، $\mu(s)$ ، را با استفاده از میانه‌پالایش و اسپلاین برآورد کردیم، سپس مقادیر برآورد شده‌ی روند را با $\delta(s)$ ، مقادیر پیشگویی شده با استفاده از کریگینگ معمولی روی باقی‌مانده‌های حاصل از دو روش فوق جمع شدند. در نهایت مقادیر به‌دست آمده با مقادیر واقعی داده‌های مجموعه‌ی آزمون UCS مقایسه شدند. برای بررسی کارکرد پیشگویی سه بعدی هر سه روش میانه‌پالایش، اسپلاین و کریگینگ عام از مقدار عددی و نمودار جعبه‌ای MSPE بهره بردیم.

نمودارهای پراکنش داده‌ها بر حسب سه محور مختصات در شکل ۷.۳ نشان داده شده‌اند.



شکل ۷.۳: نمودار پراکنش داده‌ها نسبت به مختصات طول، عرض و عمق به ترتیب (آ)، (ب) و (ج).

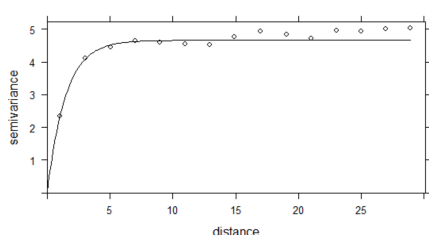
شکل ۸.۳ نمودار نیم‌تغییرنگار روی باقی‌مانده‌های حاصل از سه روش را نمایش می‌دهد، شایان ذکر است که نمودار نیم‌تغییرنگار تجربی و مدل منتخب برازش داده‌شده روی داده‌های اصلی در شکل ۵.۲ (آ)، نمایش داده شده، همچنین مقادیر اترقطعه‌ای، ازاره و دامنه‌ی

نیم‌تغییرنگار تجربی روی باقی‌مانده‌های به‌دست آمده از روش‌های مذکور در جدول ۳.۳ ارائه شده است.

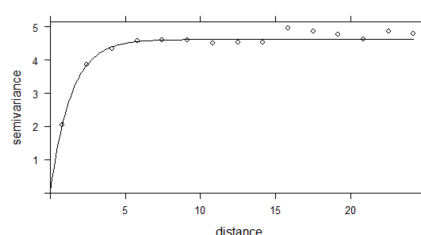
جدول ۳.۳: پارامترهای تغییرنگار

کرینگینگ عام	کرینگینگ میانه‌پالایش	کرینگینگ اسپلین	
ازاره	۰.۵۹۷	۴.۵۸۱	
دامنه	۰.۶۱۱	۱.۳۶۴	
اثر قطعه‌ای	۰.۴۲۳	۰.۰۳۵	

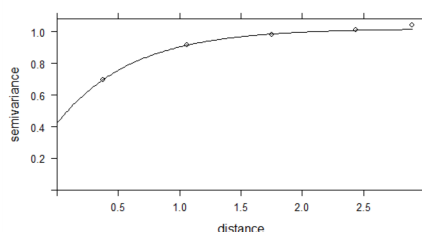
بر اساس مقادیر MSPE که در جدول ۴.۳ ارائه شده، می‌توان نتیجه گرفت که پیشگویی فضایی با روش اسپلین از MSPE کوچک‌تری نسبت به دو رقیب خود برخوردار است، پس این روش نسبت به دو روش دیگر در پیش‌گویی فضایی برتری دارد. در نهایت می‌توان نتیجه گرفت که روش کرینگینگ عام نسبت به میانه‌پالایش برتری دارد. نمودارهای جعبه‌ای شکل ۹.۳ نیز این موضوع را تایید می‌کنند. دلیل ضعیف عمل کردن میانه‌پالایش با توجه به شکل ۷.۳ را می‌توان خطی بودن روند داده‌ها نسبت به سه مختصات طول، عرض و عمق دانست، که در این‌صورت برآورد روند با رگرسیون که بر حسب میانگین عمل می‌کند کاراتر از میانه‌پالایش است که بر حسب میانه عمل می‌کند.



(ب)



(آ)

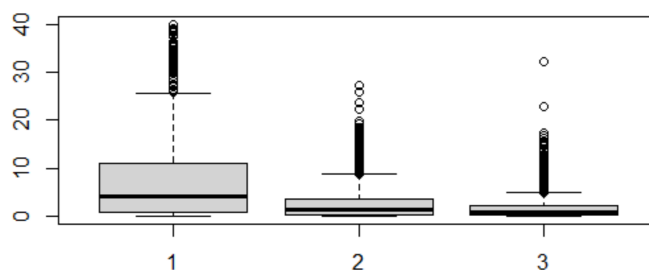


(ج)

شکل ۸.۳: نمودار سه بعدی نیم‌تغییرنگار و منحنی مدل پارامتری نمایی برازش داده شده به باقی‌مانده‌های حاصل از روش‌های (آ): اسپلین، (ب): ماکسیمم درست‌نمایی و (ج): میانه‌پالایش.

جدول ۴.۳: پارامترهای تغییرنگار

مقادیر MSPE	کریگینگ عام	کریگینگ میانه پالایش	کریگینگ اسپلین
۲.۷۹	۹.۸۹	۱.۷۸	



شکل ۹.۳: نمودار جعبه‌ای مقادیر MSPE حاصل از روش‌های ۱: کریگینگ میانه پالایش، ۲: کریگینگ عام و ۳: کریگینگ معمولی با استفاده از اسپلین.

فصل ۴

رهیافت بیزی تقریبی برای مدل‌بندی فضایی سه بعدی

۱.۴ مقدمه

آمار بیزی یک روش عقلانی در به‌کارگیری عدم قطعیت حالت‌های طبیعت به بیان احتمال است. در چارچوب فکری آمار بیزی، میزان باورها به حالت طبیعت مبنای کار است. این میزان از باورها با انتخاب باورهای پیشین قبل از مشاهده‌ی داده آغاز و با به‌کار گرفتن داده‌های به‌دست آمده از مشاهده یا آزمایش به روز می‌شوند و از آن با عنوان باورهای پسین یاد می‌کنیم. در روش‌های بیزی تحلیل بر باورهای پسین استوار است (گلمن و همکاران، ۲۰۱۳). در استنباط مبتنی بر درست‌نمایی، پارامترها^۱ ثابت در نظر گرفته می‌شوند که مقادیر آن‌ها ناشناخته است و به‌طور مستقیم نمی‌توان مقدار آن را اندازه‌گیری کرد. در استنباط بیزی، پارامترها متغیرهای تصادفی بدون مشاهده (و غیرقابل مشاهده) در نظر گرفته می‌شوند (دیگل و جورجی، ۲۰۱۹).

روش‌های MCMC به‌طور گسترده برای استنباط بیزی استفاده می‌شوند، اما محدودیت آن‌ها در بار محاسباتی‌اشان است. پیشرفت‌های موجود در جمع‌آوری داده‌ها، منجر به در

^۱ترجمه تحت اللفظی یونان باستان از پارامتر اصطلاح "Beyond measurement" به معنای "فراتر از اندازه‌گیری" است.

دسترس بودن مجموعه داده‌های بزرگ^۲ شده است که با موقعیت‌های تعیین شده‌ی فضایی، فضایی-زمانی و همچنین با منابع مختلف مشخص می‌شوند. پیچیدگی مدل برای در نظر گرفتن ساختارهای مکانی داده‌های بزرگ می‌تواند منجر به چندین روز زمان محاسبات برای انجام استنباط بیزی از طریق الگوریتم MCMC شود. به‌ویژه زمانی که بعد میدان پنهان^۳ بزرگ باشد، همچنین وجود همبستگی بین عناصر میدان پنهان و ویژگی‌های همگرایی و درهم‌تنیدگی، این الگوریتم‌ها را با مشکل روبرو می‌کنند. برای غلبه بر این چالش، رو و همکاران در سال ۲۰۰۹ در مقاله‌ی خود یک روش با نام تقریب لاپلاس آشیانه‌ای جمع‌بسته ارائه دادند. آن‌ها با استفاده از این الگوریتم ثابت کردند که می‌توانند به جواب‌های سریع و دقیقی دست یابند. الگوریتم روش INLA ابتدا به‌صورت یک برنامه مستقل شروع به کار کرد، سپس در R (به عنوان بسته‌ای با نام R-INLA) مورد دسترسی قرار گرفت و از آن پس، بین آماردان‌ها و محققان به‌ویژه در زمینه آمار فضایی از محبوبیت زیادی برخوردار شد. سایت <http://www.r-inla.org> یک سایت عالی و فعال در زمینه‌ی روش INLA است که مقالات و آموزش‌ها را در اختیار کاربران قرار می‌دهد و در تالار گفتگو محققان می‌توانند سوالات و مشکلات خود را برای بررسی و پاسخگویی، مطرح کنند (بلانگیادو و کملتی، ۲۰۱۵).

طی چندین سال، محققانی همچون گیلز و همکاران (۱۹۹۶) و بروکس و همکاران (۲۰۱۱) از رویکرد MCMC برای محاسبه‌ی توزیع پسین توأم پارامترهای مدل استفاده می‌کردند. اما این روش بار محاسباتی بسیار بالایی دارد چرا که توزیع در ابعاد بالاتر محاسبه می‌شود.

اولین بار رو و همکاران (۲۰۰۹) INLA را پیشنهاد دادند که در استنباط بیزی سریع‌تر عمل می‌کند. آن‌ها متذکر شدند که به‌جای برآورد توزیع پسین توأم پارامترهای مدل، می‌توان روی توزیع‌های پسین کناری پارامترهای مدل متمرکز شد. در بسیاری از موارد، استنباط کناری برای دانش از پارامترهای مدل و اثرات پنهان کافی است و نیازی به توزیع چندمتغیره پسین توأم نیست. در ادامه آن‌ها توجه خود را به مدل‌هایی معطوف کردند که می‌توانند به‌صورت میدان‌های گوسی تصادفی مارکوفی^۴ (GMRF) بیان شوند (گومز، ۲۰۲۰). این موضوع زمان محاسبات مدل را کاهش می‌دهد (رو و هلد، ۲۰۰۵). INLA در رده گسترده‌ای از مدل‌ها به نام مدل‌های گوسی پنهان^۵ (LGMs) کاربرد دارد. عضویت در این رده مستلزم این است که برخی از پارامترهای موجود در مدل دارای توزیع گوسی باشند (وانگ و همکاران، ۲۰۱۸). به دنبال پیشنهاد این رویکرد جدید، مارتینز و همکاران (۲۰۱۳) و اخیراً رو و همکاران (۲۰۱۷) مقالاتی را در این رابطه ارائه داده‌اند.

^۲ Big datasets

^۳ Latent field

^۴ Gaussian Markov random field

^۵ Latent Gaussian Models

۲.۴ مدل‌های گوسی پنهان

مدل‌های گوسی پنهان زیرمجموعه‌ای از مدل‌های بیزی سلسله‌مراتبی هستند. در این نوع از مدل‌ها، پارامتر در سطح نهفته و دارای توزیع پیشین توأم به شرط ابرپارامترها است. رده LGMها یک زیررده از مدل‌های رگرسیونی با ساختار جمعی هستند. رو و همکاران در پژوهش خود فرض کردند توزیع مشاهده‌ها y_i ، برای $i = 1, \dots, n$ ، از یک خانواده و اغلب از خانواده توزیع‌های نمایی تبعیت می‌کنند. میانگین یا یک چندک مشخص مثل μ_i ، و پیشگوی جمعی مدل رگرسیونی η_i از طریق تابع پیوند $g(\cdot)$ به شکل $g(\mu_i) = \eta_i$ مدل‌بندی می‌شود. در واقع

$$\eta_i = \alpha + \sum_{j=1}^J f_j(\mathbf{u}_{i,j}) + \sum_{k=1}^K \beta_k z_{i,k} + \varepsilon_i \quad (1.4)$$

که در آن α عرض از مبدا، $\{f_j\}_{j=1}^J$ اثرات ناپارامتری نامعلوم با برخی وابستگی‌های ساختاری خاص با وزن‌های مشخص $\mathbf{u}$ و ضرایب $\{\beta_k\}_{k=1}^K$ اثرات ثابت خطی متغیرهای تبیینی $z_{i,1}, \dots, z_{i,K}$ روی متغیر پاسخ است و $\varepsilon_1, \dots, \varepsilon_n$ خطاهای غیرساختاری مدل هستند. **گومز** در کتاب خود در سال ۲۰۲۰، اشاره کرده است که تابع درستنمایی مشاهدات y_i ها، باید باهم مرتبط باشند و الزامی در اینکه از خانواده‌ی توزیع‌های نمایی باشند وجود ندارد. این رده از مدل‌ها، چون تابع $\{f_j\}$ ساختار نامعلومی دارد می‌تواند کاربردهای زیادی داشته باشد. مدل‌های گوسی پنهان، یک زیرمجموعه از تمام مدل‌های جمعی بیزی با ساختار پیشگوی جمعی است، یعنی پیشین گوسی به α ، $\{\beta_k\}$ ، $\{f_j\}$ و $\{\varepsilon_i\}$ اختصاص داده می‌شود. در چارچوب INLA بردار اثرات پنهان $\mathbf{x}$ به صورت زیر نمایش داده می‌شود که دارای ابرپارامترهای پنهان است:

$$\mathbf{x} = \{\eta_i, \alpha, \beta_k, f_j\}.$$

پارامترها در $\mathbf{x}$ دارای توزیع توأم گوسی هستند که روی ابرپارامترهای $\theta = (\theta_1, \dots, \theta_p)^T$ شرطی شده است. توزیع پسین توأم اثرات پنهان $\mathbf{x}$ و ابرپارامترهای θ را می‌توان به صورت زیر نوشت:

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

در معادله‌ی قبلی، $\pi(\mathbf{y})$ درستنمایی کناری مدل را نشان می‌دهد که یک ثابت نرمال‌ساز است و اغلب نادیده گرفته می‌شود، چرا که محاسبه‌ی آن معمولاً دشوار است. با این حال INLA، یک تقریب دقیق از آن را ارائه می‌دهد (**هوبین و استورویک**، ۲۰۱۶). ابرپارامترها معمولاً دارای چگالی پسین گوسی نیستند و اغلب واریانس کناری، دامنه همبستگی و همواری یا سایر پارامترهای همبستگی اثرات تصادفی را مشخص می‌کنند. به صورت شماتیک LGMها با تابع پیوند یک متغیره ممکن است از نظر سطح داده، سطح پنهان و سطح ابرپارامترها به گونه‌ی سلسله‌مراتبی نشان داده شوند (**رانکسن و همکاران**، ۲۰۲۰).

سطوح داده: فرض کنید مشاهدات $\mathbf{y} = (y_1, \dots, y_n)^T$ به پارامترهای پنهان $\mathbf{x}$ وابسته و دارای چگالی

است. $\pi(y|x, \theta)$ برای سادگی فرض می‌شود با استقلال شرطی، توزیع n مشاهده به وسیله‌ی درستنمایی زیر محاسبه می‌شود:

$$\pi(y|x, \theta) = \prod_{i=1}^n \pi(y_i|x, \theta) \quad (2.4)$$

به‌طوری که $\pi(y_i|x, \theta) = \pi(y_i|\eta_i, \theta)$ با $\mu_i = Q_{\pi}(y_i|\eta_i) = g^{-1}(\eta_i)$ که در آن μ_i میانگین یا یک چندک مناسب از تابع چندک Q_p است. در رابطه‌ی (۲.۴) هر عضو از y_i تنها به یک عضو از x_i از میدان پنهان x مرتبط است.

سطوح پنهان: پارامتر پنهان x ، دارای چگالی پیشین گوسی است و به‌طور بالقوه به برخی از ابرپارامترها وابسته است. تابع چگالی x ممکن است به‌صورت

$$\pi(x|\theta) = N(x|\mu(\theta), \Sigma(\theta))$$

نوشته شود. در رابطه‌ی فوق، عبارات سمت راست چگالی گوسی با بعد برابر بعد x با بردار میانگین $\mu(\theta)$ و ماتریس کوواریانس $\Sigma(\theta)$ را نشان می‌دهد (**گلمن و همکاران**، ۲۰۱۳).

سطوح ابرپارامترها: ابرپارامترهای θ ، دارای چگالی پیشین $\pi(\theta)$ هستند.

فرض می‌کنیم x دارای توزیع چندمتغیره گوسی با بردار میانگین μ و ماتریس دقت $Q(\theta) = \Sigma(\theta)^{-1}$ باشد، یعنی

$$\pi(x|\theta) = (\pi)^{-\frac{n}{2}} |Q(\theta)|^{\frac{1}{2}} \exp\left(-\frac{1}{2} x^T Q(\theta) x\right) \quad (3.4)$$

به‌طوری که مولفه‌های میدان گوسی x به‌طور شرطی مستقل هستند. از این موضوع می‌توان نتیجه گرفت که ماتریس دقت $Q(\theta)$ ، **تُنک**^۶ است. **رو و هلد** (۲۰۰۵) این خصوصیت را GMRF تعریف کردند. زمانی که استنباط با استفاده از GMRF صورت گیرد، **تُنک** بودن ماتریس دقت به کم شدن زمان محاسبات می‌انجامد.

توزیع پسین x و θ از ضرب درستنمایی (۲.۴) و چگالی GMRF (۳.۴) و توزیع پیشین ابرپارامترها به‌صورت زیر محاسبه می‌شود:

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

^۶ Sparse

۳.۴ تقریب لاپلاس آشیانه‌ای جمع بسته

هدف استنباط بیزی، تقریب پسین کناری اعضای میدان گوسی پنهان x

$$\pi(x_i|y) = \int \pi(x_i, \theta|y) d\theta = \int \pi(x_i|\theta, y) \pi(\theta|y) d\theta \quad (۴.۴)$$

و اعضای بردار ابرپارامتر

$$\pi(\theta_j|y) = \int \pi(\theta|y) d\theta_{-j} \quad (۵.۴)$$

است، که در آن θ_{-j} بردار ابرپارامترها در غیاب عضو j ام θ_j است. این تقریب در سه گام انجام می‌شود.

گام اول: ابتدا تقریب پسین کناری θ با استفاده از تقریب لاپلاس محاسبه می‌شود. رو و همکاران (۲۰۰۹) معتقدند اساس رویکرد INLA از تقریب پسین کناری θ یعنی $\tilde{\pi}(\theta|y)$ پیروی می‌کند.

$$\tilde{\pi}(\theta|y) \propto \frac{\pi(x, \theta, y)}{\tilde{\pi}_G(x|\theta, y)} \Big|_{x=x^*(\theta)} \quad (۶.۴)$$

که در آن، $\tilde{\pi}_G(x|\theta, y)$ تقریب گوسی برای چگالی شرطی کامل x و $x^*(\theta)$ مقدار مد شرطی x برای مقادیر داده شده θ است. تقریب چگالی (شرطی) θ_j که $j = 1, \dots, J$ به صورت

$$\tilde{\pi}(\theta_j|y) = \int \tilde{\pi}(\theta|y) d\theta_{-j}$$

است.

گام دوم: تقریب لاپلاس پسین کناری اعضای میدان x را محاسبه می‌کنیم. به طور ساده‌تر، تقریب لاپلاس $\tilde{\pi}(x_i|\theta, y)$ برای θ های تعیین شده به صورت زیر محاسبه می‌شود:

$$\tilde{\pi}(x_i|y) = \int \tilde{\pi}(x_i, \theta|y) d\theta = \int \tilde{\pi}(x_i|\theta, y) \pi(\theta|y) d\theta.$$

گام سوم: در آخرین گام، دو مرحله‌ی قبل را با استفاده از انتگرال گیری عددی با هم ترکیب کرده و انتگرال توزیع پسین میدان پنهان را با استفاده از جمع متناهی به دست می‌آوریم. در واقع داریم

$$\tilde{\pi}(x_i|\theta, y) = N\{x_i; \mu_i(\theta), \sigma_i^2(\theta)\}$$

که در آن $\mu(\theta)$ (بردار) میانگین تقریب گوسی و $\sigma^2(\theta)$ (بردار) واریانس کناری متناظر است. در نهایت

$$\tilde{\pi}(x_i|y) = \sum_{k=1}^n \tilde{\pi}(x_i|\theta_k, y) \tilde{\pi}(\theta_k|y) \Delta_k \quad (۷.۴)$$

که در آن Δ_k وزن‌هایی هستند که به θ_k نسبت داده می‌شود. رو و مارتینونو (۲۰۰۷) نشان دادند که تقریب کناری پسین θ دقیق است.

۱.۳.۴ تقریب چگالی پسین توأم بردار ابرپارامترها

چگالی پسین توأم بردار ابرپارامترها به صورت

$$\begin{aligned}\pi(\boldsymbol{\theta}|\mathbf{y}) &= \frac{\pi(\boldsymbol{\theta}, \mathbf{y})}{\pi(\mathbf{y})} \\ &= \frac{\pi(\boldsymbol{\theta}, \mathbf{y}) \pi(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y})}{\pi(\mathbf{y}) \pi(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y})} \\ &= \frac{\pi(\mathbf{x}, \boldsymbol{\theta}, \mathbf{y})}{\pi(\mathbf{y}) \pi(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y})} \\ &= \frac{\pi(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}) \pi(\mathbf{x}|\boldsymbol{\theta}) \pi(\boldsymbol{\theta})}{\pi(\mathbf{y}) \pi(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y})} \\ &\propto \frac{\pi(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}) \pi(\mathbf{x}|\boldsymbol{\theta}) \pi(\boldsymbol{\theta})}{\pi(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y})}\end{aligned}\quad (۸.۴)$$

محاسبه می‌شود. تقریب لاپلاس آشیانه‌ای جمع بسته‌ی معادله‌ی (۸.۴) برای هر مقدار $\boldsymbol{\theta}$ که پیش‌تر در معادله‌ی (۶.۴) به آن اشاره شد به شکل

$$\tilde{\pi}(\boldsymbol{\theta}|\mathbf{y}) \propto \frac{\pi(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}) \pi(\mathbf{x}|\boldsymbol{\theta}) \pi(\boldsymbol{\theta})}{\pi(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y})} \bigg|_{\mathbf{x}=\mathbf{x}^*(\boldsymbol{\theta})} \quad (۹.۴)$$

به دست می‌آید. در معادله‌ی (۹.۴) مخرج کسر از توزیع گوسی حاصل می‌شود. با استفاده از بسط تیلور در تقریب گوسی، توزیع شرطی کامل

$$\pi(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y}) \propto \left(-\frac{1}{2} \mathbf{x}^T \mathbf{Q} \mathbf{x} + \sum_{i \in \mathcal{I}} \log \pi(y_i | x_i, \boldsymbol{\theta}) \right)$$

حول مُد توزیع شرطی کامل $\mathbf{x}$ ، یعنی $\mathbf{x}^*(\boldsymbol{\theta}) = (x_1^*(\boldsymbol{\theta}), \dots, x_m^*(\boldsymbol{\theta}))$ به صورت زیر محاسبه می‌شود:

$$\begin{aligned}\tilde{\pi}_G(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y}) &\propto \left(-\frac{1}{2} \mathbf{x}^T \mathbf{Q} \mathbf{x} + \sum_{i \in \mathcal{I}} g_i(x_i) \right) \\ &\propto \left(-\frac{1}{2} \mathbf{x}^T \mathbf{Q} \mathbf{x} + \sum_{i \in \mathcal{I}} (a_i + b_i x_i - \frac{1}{2} c_i x_i^2) \right) \\ &\propto \left(-\frac{1}{2} \mathbf{x}^T (\mathbf{Q} + \text{diag}(\mathbf{c})) \mathbf{x} + \mathbf{b}^T \mathbf{x} \right)\end{aligned}\quad (۱۰.۴)$$

که در آن $a_i = g'(x_i^*(\boldsymbol{\theta})) + x_i^*(\boldsymbol{\theta}) c_i$ ، $b_i = g'(x_i^*(\boldsymbol{\theta}))$ و $c_i = -g''(x_i^*(\boldsymbol{\theta}))$ (۱۰.۴) مُد با استفاده از روش نیتون-رافسون^۷ که با نام الگوریتم امتیازدهی فیشر^۸ نیز شناخته می‌شود، محاسبه می‌شود (فارمر و تاتز، ۲۰۰۱). اگر $\mu^{(\circ)}$ حدس اولیه باشد، $g_i(x_i)$ حول $\mu_i^{(\circ)}$

^۷Newton-Raphson method

^۸Fisher scoring algorithm

به مرحله دوم بسط داده می‌شود و می‌توان نوشت

$$g_i(x_i) \approx g_i(\mu_i^{(0)}) + b_i x_i - \frac{1}{2} c_i x_i^2$$

که $\{b_i\}$ و $\{c_i\}$ به $\mu^{(0)}$ وابسته هستند. این یک تقریب گوسی بهینه با ماتریس دقت $Q + \text{diag}(c)$ است و مُد از حل رابطه‌ی $Q + \text{diag}(c)\mu^{(1)} = b$ به‌دست می‌آید. این فرآیند تا زمانی که به یک توزیع پنهان با میانگین x^* و ماتریس دقت $Q^* = Q + \text{diag}(c^*)$ همگرا شود، تکرار می‌شود (رو و همکاران، ۲۰۰۹).

تقریب لاپلاس چگالی پسین $\pi(\theta|y) = \frac{\pi(x, \theta, y)}{\pi(x|\theta, y)}$ با بسط مخرج کسر حول $x = x^*(\theta)$ به‌صورت

$$\tilde{\pi}_{LA}(\theta|y) = \left. \frac{\pi(x, \theta, y)}{\pi_G(x|\theta, y)} \right|_{x=x^*(\theta)}$$

به‌دست می‌آید.

در سه مرحله از فرآیند اجرای روش INLA، از تقریب $\tilde{\pi}_{LA}(\theta|y)$ استفاده می‌شود. مهمترین و اولین کاربرد آن در به‌دست آوردن تقریب بیزی پسین میدان پنهان است. در مرحله‌ی بعد از رابطه‌ی (۹.۴) برای محاسبه‌ی تقریب کناری درست‌نمایی $\pi(y)$ استفاده می‌شود. در نهایت از این تقریب در جهت محاسبه‌ی چگالی کناری پسین بعضی از اعضای ابرپارامترهای θ ، یعنی $\pi(\theta_i|y)$ استفاده می‌شود. باید توجه داشت که با انتگرال‌گیری روی محدوده‌ی چند بعدی θ می‌توان از (۹.۴) برای محاسبه‌ی سه مرحله‌ی مذکور اقدام کرد. پس ضرورت انتخاب نقاط مناسب انتگرال‌گیری کاملاً مشهود است. رو و همکاران (۲۰۰۹) دو رویکرد توری^۹ و طرح مرکب مرکزی^{۱۰} (CCD) را برای انتخاب نقاط انتگرال‌گیری مناسب پیشنهاد دادند. رویکرد توری، مشتمل از نقاط یک توری است که $\tilde{\pi}_{LA}(\theta|y)$ چگالی بیشتری در آن نقاط دارد. زمانی که بعد θ بزرگ نباشد (مثلاً کمتر از ۶)، این روش بسیار دقیق و مناسب است. وقتی بعد θ بین ۶ تا ۱۲ باشد، استفاده از رویکرد CCD مفید است. در هر دو رویکرد، $\tilde{\pi}_{LA}(\theta|y)$ تک مُدی فرض می‌شود. در فرآیند انتگرال‌گیری به مُد $\tilde{\pi}_{LA}(\theta|y)$ نیاز است و با نماد θ^* نشان داده می‌شود. با توجه به پیچیدگی تابع مورد نظر می‌توان از انتگرال‌گیری بهینه‌سازی نیوتن-رافسون یا شبه‌نیوتن^{۱۱} برای بهینه کردن θ^* استفاده کرد.

۲.۳.۴ روش توری

روش توری طی سه گام انجام می‌شود که در ادامه به بررسی آن‌ها می‌پردازیم.

گام اول: مقدار $\log\{\tilde{\pi}(\theta|y)\}$ را نسبت به θ بهینه می‌کنیم تا مقدار مُد $\tilde{\pi}(\theta|y)$ محاسبه شود. این کار را می‌توان با روش نیوتن-رافسون انجام داد.

^۹Grid

^{۱۰}Central composite design

^{۱۱}Quasi-Newton

گام دوم: ماتریس منفی هسین با استفاده از تفاضل‌های متناهی در پیکره‌ی θ^* محاسبه می‌شود. با فرض $\Sigma = H^{-1}$ ، اگر چگالی گوسی را داشته باشیم، آنگاه Σ ماتریس کوواریانس θ خواهد بود. اگر ماتریس Σ به صورت $\Sigma = V\Lambda V$ تجزیه شود، برای ساده‌سازی متغیر استاندارد z به جای θ به صورت

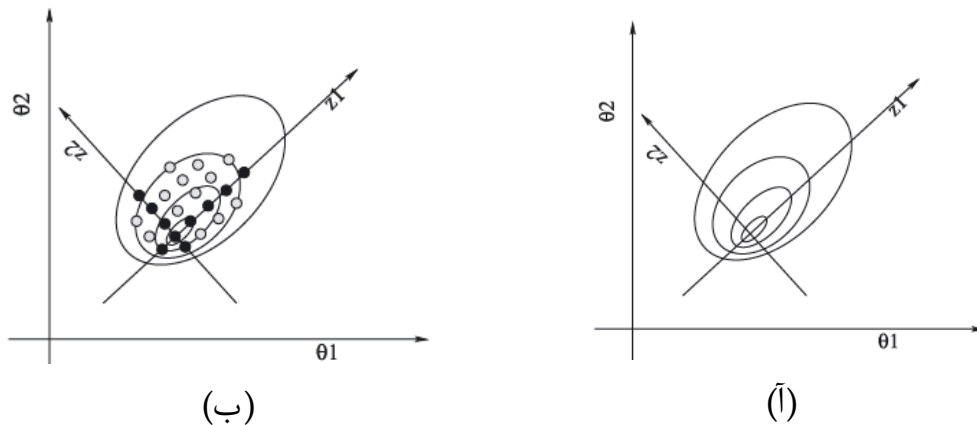
$$\theta(z) = \theta^* + V\Lambda^{1/2}z$$

استفاده می‌شود.

گام سوم: بررسی $\log\{\tilde{\pi}(\theta|y)\}$ با استفاده از متغیر z . برای توضیح ساده و قابل درک، شکل ۱.۴ نمودار تراز $\log\{\tilde{\pi}(\theta|y)\}$ را وقتی $m = 2$ است برای مُد و مختصات جدید نشان می‌دهد. در واقع هدف کاوش $\log\{\tilde{\pi}(\theta|y)\}$ برای یافتن بیشترین چگالی احتمال است. برای این کار از نقطه‌ی مُد $z = 0$ شروع می‌کنیم. سپس یک گام در راستای مثبت z_1 به طول δ_z برمی‌داریم تا جایی که شرط

$$\log[\{\tilde{\pi}(\theta(0)|y)\}] - \log[\{\tilde{\pi}(\theta(z)|y)\}] < \delta_\pi \quad (11.4)$$

برقرار شود. برای z_2 نیز به همین ترتیب رفتار می‌کنیم. با این کار نقاط سیاه به وجود می‌آیند. اکنون از همین روش می‌توان تمام نقاط میانی (نقاط خاکستری) را با ترکیبات مختلف نقاط سیاه ایجاد کرد و مختصات جدید تولید می‌شوند. رو و همکاران در پژوهش خود، $\delta_z = 1$ و $\delta_\pi = 2/5$ را در نظر گرفتند (رو و همکاران، ۲۰۰۹).



شکل ۱.۴: (آ) مکان مُد، و سیستم مختصات z ، (ب) نمایش نقاط منتخب انتگرال‌گیری با راهبرد توری

۳.۳.۴ روش طرح مرکب مرکزی

انتخاب نقاط انتگرال‌گیری در روش توری، وقتی تعداد ابرپارامترها، m ، زیاد می‌شود خیلی پرهزینه است، حتی اگر تعداد ابرپارامترها متوسط (۶ تا ۱۲) باشد. مثلاً اگر چگالی $\tilde{\pi}(\theta|z)$

گوسی باشد و $\delta_z = 1$ و $\delta_\pi = 2/5$ آن گاه به $\mathcal{O}(\delta^m)$ نقطه برای به‌دست آوردن (۷.۴) نیاز است و اگر این روش به سه نقطه در هر بعد محدود شود به $\mathcal{O}(3^m)$ نقطه نیاز است که از مرتبه‌ی نمایی است (رو و همکاران، ۲۰۰۹).

در این بخش رویکرد CCD شرح داده می‌شود که هزینه‌ی محاسبات را برای m های بزرگ کاهش می‌دهد. روش CCD به عنوان مساله‌ی طرح رویه‌ی پاسخ در نظر گرفته می‌شود و نقاط روی فضای m بعدی پخش می‌شوند و بر اساس پاسخ‌های به‌دست آمده در آن نقاط خمیدگی پسینی حول مُد برآورد می‌شود. در این بخش، تنها سطوح پاسخ از مرتبه دو را در نظر گرفته و از طرح درجه دو CCD استفاده می‌شود (باکس و ویلسون، ۱۹۵۱). در اغلب زمینه‌های پژوهشی مختلف، پژوهشگران برای کشف اطلاعاتی درباره‌ی فرآیند سیستم‌های خاص، آزمایش‌هایی را انجام می‌دهند. آزمایش‌گر برای تشخیص عوامل تغییرات در پاسخ، باید عوامل متغیر در آزمایش، دامنه‌ی تغییرات هر عامل و سطوح خاصی که در هر بار اجرا باید مورد استفاده قرار گیرند را تعیین کند. طرح یک آزمایش، ابزار مهمی برای تعیین عملکرد درست فرآیند تولید است. در علم آمار منظور از طرح آزمایش آماری این است که فرآیندی را طراحی کنند که بتوان با آن داده‌های مناسبی به روش‌های مختلف آماری جمع‌آوری و تحلیل کرد.

طرح 2^k عاملی، طرح‌های k عاملی هستند که هر عامل فقط از دو سطح کمی یا کیفی برخوردار است. طرح عاملی V -تجزیه^{۱۲}، به طرح‌های عاملی می‌گویند که اثرهای اصلی و اثرهای متقابل دو عاملی با اثرهای اصلی و اثرهای متقابل در عامل دیگری با یکدیگر، هم‌اثر نباشند. اگر تعداد عامل‌ها در این نوع طرح‌ها زیاد باشد، اجرا از حیث تعداد و تعیین نقاط طرح کار دشواری است. بزرگترین طرح عاملی V -تجزیه، 3^{10-3} برابر با 4^{11-4} است (مونتگومری، ۲۰۰۰؛ باکس و همکاران، ۱۹۷۸). سانچز و سانچز (۲۰۰۵)، با استفاده از ماتریس مرتب‌سازی شده‌ی هادامارد^{۱۳} (بیوچمپ، ۱۹۸۴) و اندیس والش^{۱۴}، راهکار یادگیری را برای تعیین طرح‌های عاملی کسری^{۱۵}، حتی تا 12^0 عامل، ارائه دادند که به‌صورت مختصر در ادامه بیان می‌شود. ماتریس هادامارد H_k ، یک ماتریس $2^k \times 2^k$ با درایه‌های $+1$ یا -1 است که می‌توان آن را با روابط بازگشتی به‌صورت

$$H_1 = [1], \quad H_2 = \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}, \quad H_{k+1} = \begin{bmatrix} H_k & H_k \\ H_k & -H_k \end{bmatrix}, \quad (12.4)$$

محاسبه کرد. اگر ستون‌های H_k با $h_0, h_1, \dots, h_{N-1}$ تعیین شوند که در آن $N = 2$ ، هر ستون h_i به اثر اصلی یا اثر متقابل دو عاملی با استفاده از الگوریتم بازگشتی برای اندیس‌گذاری والش

^{۱۲}V-Resoulution

^{۱۳}Hadamard ordered matrix

^{۱۴}Walsh index

^{۱۵}Fractional factorial design

طرح عاملی V -تجزیه کسری تجزیه می‌شود. در k مرحله این اندیس‌گذاری انجام می‌شود که در مرحله i ام آن، اندیس عامل x_i تعیین می‌شود. اگر اندیس والش عامل‌های x_{j1} و x_{j2} به ترتیب a_1 و a_2 باشند، اندیس اثر متقابل بین این دو عامل، یعنی اثر $x_{j1}x_{j2}$ ، از رابطه‌ی $a_1 \oplus a_2$ محاسبه می‌شود، که در آن بُعد از تغییر مبنای اعداد به دو، یای انحصاری $\oplus$ ، با استفاده از جدول ۱.۴ محاسبه می‌شود.

جدول ۱.۴: نقاط طرح هادامارد والش برای $k = 3$

A	B	$A \oplus B$
۰	۰	۰
۱	۰	۱
۰	۱	۱
۱	۱	۰

به عنوان مثال $7 \oplus 4 = (111)_2 \oplus (100)_2 = (011)_2 = (1 \times 2^0) + (1 \times 2^1) + (0 \times 2^2) = 3$. در انتهای مرحله‌ی i ام، $i = 1, \dots, k$ دو مجموعه‌ی A_i و T_i به صورت‌های

$$A_i = \{X_1, \dots, X_k \text{ اثرهای اصلی عامل‌های}\}$$

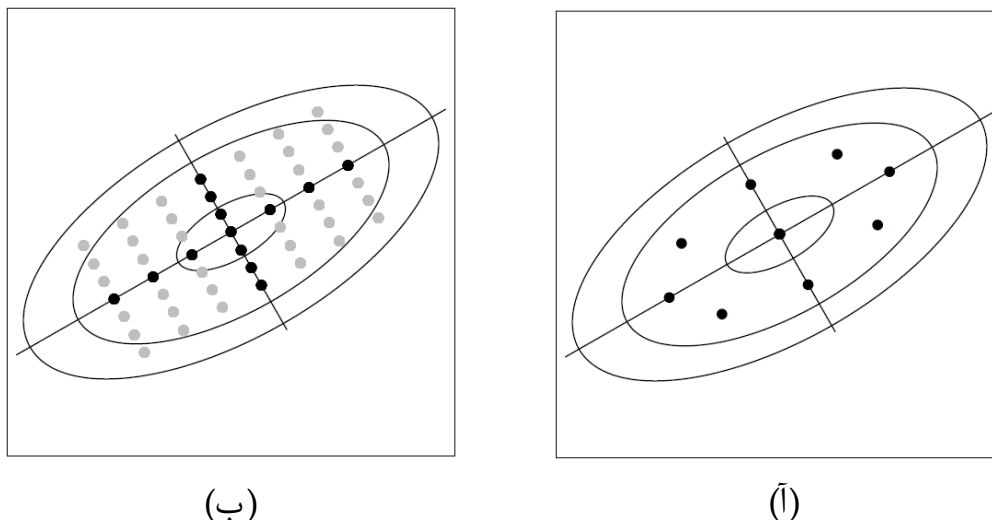
$$T_i = \{A_i \text{ اثرهای متقابل دو عاملی بین مولفه‌های اندیس والش متناظر}\}$$

به دست می‌آید که A_0 و T_0 تهی فرض می‌شوند. برای اعتبار اندیس والش، اندیس پیشنهاد شده‌ی عامل X_j که با b نشان داده می‌شود، باید دارای دو شرط زیر باشد:

۱. نباید متعلق به A_{j-1} و T_{j-1} باشد.

۲. اندیس اثرهای متقابل دو عاملی بین x_j و $x_1, \dots, x_{j-1}$ نباید متعلق به A_j و T_j باشند.

در شکل ۲.۴ نقاط منتخب انتگرال‌گیری به روش‌های CCD و روش توری نشان داده شده‌اند.



شکل ۲.۴: نقاط منتخب انتگرال‌گیری با استفاده از (آ) روش CCD و (ب) روش توری

الگوریتم تعیین اندیس والش در ادامه ذکر شده‌است.

الگوریتم ۱.۳.۴. تعیین اندیس والش

۱. قرار می‌دهیم $m = 1$.

۲. $b = \max\{i : i \in A_{m-1}\}$.

۳. تا وقتی $b \in T_{m-1}$ یا $\{ \exists j \in A_{m-1} : b \oplus i \in A_{m-1} \cup T_{m-1} \}$ قرار می‌دهیم $b = b + 1$.

۴. دو مجموعه‌ی $A_m = \{b\} \cup A_{m-1}$ و $T_m = T_{m-1} \cup \{i \oplus b : i \in A_m\}$ را محاسبه می‌کنیم.

۵. قرار می‌دهیم $m = m + 1$ و تا وقتی $m < k$ به گام ۲ برمی‌گردیم.

بعد از اندیس‌گذاری به روش والش، هر ستون ماتریس هادامارد متعلق به یک اثر اصلی یا اثر متقابل دو عاملی خواهد بود و هر سطر ماتریس مشخص‌کننده‌ی سطوح عامل در یک اجرا خواهد بود که به آن نقاط طرح می‌گویند.

جدول ۳.۴ نقاط طرح را برای $k = 3$ نشان می‌دهد. در این طرح‌ها به تعداد نقاط غیرتکراری، طرح نیاز به اجرا دارد.

جدول ۲.۴: نقاط طرح هادامارد والش برای $k = 3$

نماد طرح	۶: $X_2 X_3$	۵: $X_1 X_3$	۴: X_3	۳: $X_1 X_2$	۲: X_2	۱: X_1
۱	+	+	+	+	+	+
۲	+	-	+	-	+	-
۳	-	+	+	-	-	+
۴	-	-	+	+	-	-
۵	-	-	-	+	+	+
۶	-	+	-	-	+	-
۷	+	-	-	-	-	+
۸	+	+	-	+	-	-

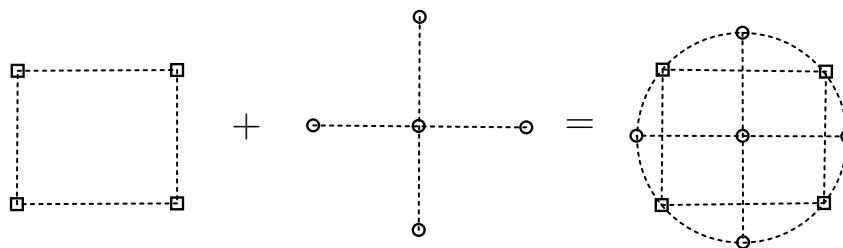
جدول ۳.۴: تعداد اجرا در طرح عاملی کسری

بعد θ	۱-۲	۳	۴-۵	۶	۷-۸	۹-۱۱	۱۲-۱۷
تعداد نقاط	۴	۸	۱۶	۳۲	۶۴	۱۲۸	۲۵۶

تعیین نقاط و تعداد نقاط در طرح مکعب مرکزی

طرح مرکب مرکزی، طرح عامل کسری را در خود دارد که با یک نقطه‌ی مرکزی و یک گروه متشکل از $2m$ نقطه با عنوان ستاره‌ای ترکیب شده است. بنابراین، تعداد k نقطه از مجموع تعداد طرح و $2m + 1$ نقطه جدید به دست می‌آیند. مثلاً در طرح CCD، تعداد نقاط مورد نیاز برای $m = 5$ برابر $27 = 1 + (2 \times 5) + 16$ است که در مقایسه با تعداد نقاط مورد نیاز در روش توری $3125 = 5^5$ یا $243 = 3^5$ بسیار اندک است.

در این روش، تمام نقاط روی سطح کره‌ی m بعدی به شعاع $\sqrt{m}$ قرار می‌گیرند. نقاط ستاره‌ای در فاصله‌ی $\pm\sqrt{m}$ روی محورها و نقطه‌ی مرکزی در مبدا قرار می‌گیرند. شکل ۳.۴ موقعیت نقاط را برای CCD وقتی $m = 2$ است را نشان می‌دهد.



شکل ۳.۴: موقعیت نقاط در طرح مرکب مرکزی

برای مشخص شدن وزن‌های انتگرال در (۷.۴) فرض شده است که $\tilde{\pi}(\theta|y)$ گوسی استاندارد است و انتگرال $\theta^T \theta$ برابر m باشد. با این شرایط فرض شده، وزن انتگرال‌گیری روی نقاط کره به شعاع $f_0 \sqrt{m}$ به صورت زیر محاسبه می‌شود:

$$\Delta_k = \left[(k-1)(f_0^2 - 1) \left\{ 1 + \exp\left(-\frac{mf_0^2}{4}\right) \right\} \right]^{-1}$$

که در آن $f_0 > 1$. وزن انتگرال‌گیری نقطه‌ی مرکزی برابر $1 - (k-1)\Delta_k$ است (باکس و ویلسون، ۱۹۵۱).

درون‌یابی گوسی نامتقارن

بعضی از چگالی‌های کناری $\pi(\theta_k|y)$ را می‌توان با تقریب توزیع توأم $\pi(\theta|y)$ با یک توزیع نرمال چند متغیره $\tilde{\pi}(\theta|y)$ محاسبه کرد. تقریب گوسی $\pi(\theta_k|y)$ محاسبات اضافی ندارد، زیرا مد θ^* و ماتریس منفی هسین H با استفاده از (۷.۴) تقریب شده است.

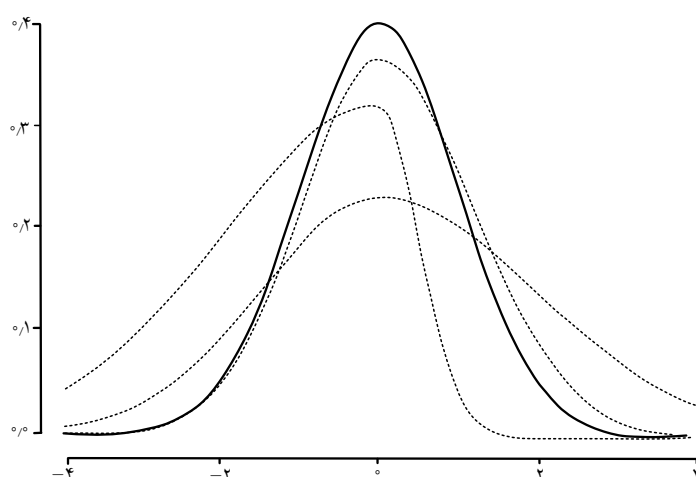
اگر $\pi(\theta_k|y)$ چوله باشد، در این صورت تقریب گوسی مناسب نیست. در این حالت، برای تصحیح تقریب گوسی با هزینه‌ی کمتر به شکل زیر عمل می‌شود. تابع $f(\theta)$ را وقتی با نقاط $z(\theta) = (z_1(\theta), \dots, z_k(\theta))$ فرض شده در z باشند، می‌توان به صورت زیر تعریف کرد:

$$f(\theta) = \prod_{k=1}^m f_k(z_k(\theta)) \quad (۱۳.۴)$$

که در آن

$$f_k(z) \propto \begin{cases} \exp\left(-\frac{1}{2(\sigma^{k+})^2} z^2\right) & z \geq 0 \\ \exp\left(-\frac{1}{2(\sigma^{k-})^2} z^2\right) & z < 0 \end{cases} \quad (۱۴.۴)$$

برای رفع نامتقارنی $\tilde{\pi}(\theta|y)$ ، پارامترهای $(\sigma^{k+}, \sigma^{k-})$ که $k = 1, \dots, m$ ، را به گونه‌ای تعریف می‌کنیم که بر پایه m محور مختلف تغییری نداشته باشد، اما بر اساس مسیر، جهت مثبت و منفی هر محور تغییر کند. برای محاسبه این توابع، ابتدا باید توجه داشته باشیم که در یک توزیع گوسی، در لگاریتم چگالی وقتی مد به ± 2 می‌رسد، انحراف معیار برابر ۲- می‌شود. این پارامترها را در روشی که به طور دقیق برای تمام مسیرها تقریب زده می‌شوند، حساب می‌کنیم. سپس تقریب‌های $\pi(\theta_k|y)$ با انتگرال‌گیری عددی (۱۳.۴) محاسبه می‌شوند، به طوری که محاسبه‌ی آن برای پارامترهای معلوم است. شکل ۴.۴ انعطاف پذیری $f_k(z)$ را در معادله‌ی (۱۴.۴) برای نقاط مختلف σ^{k+} و σ^{k-} نشان می‌دهد.



شکل ۴.۴: توزیع نرمال استاندارد (خط ممتد) و چگالی معادله‌ی (۱۴.۴) برای مقادیر مختلف پارامتر مقیاس (خطوط خط‌چین)

الگوریتم انتگرال‌گیری عددی

تقریب چگالی‌های پسین کناری $\tilde{\pi}(\theta_k | \mathbf{y})$ با استفاده از الگوریتم انتگرال‌گیری عددی به صورت زیر است:

$$\tilde{\pi}(\theta_k | \mathbf{y}) = \begin{cases} N(0, \sigma_{k+}^2) & \sigma > 0 \\ N(0, \sigma_{k-}^2) & \sigma \leq 0. \end{cases} \quad (15.4)$$

سپس این معادله برای محاسبه‌ی σ_{k+}^2 و σ_{k-}^2 بدون محاسبه با انتگرال‌گیری عددی به صورت الگوریتم قبلی، به کار می‌رود. لم زیر برای این عمل مفید است.

لم ۱.۳.۴. با فرض $\mathbf{x} = (x_1, \dots, x_n)^T \sim N(0, \Sigma)$

$$-\frac{1}{2}(x_1, E(\mathbf{x}_{-1} | x_1)^T) \Sigma^{-1} \begin{pmatrix} x_1 \\ E(\mathbf{x}_{-1} | x_1)^T \end{pmatrix} = -\frac{1}{2} \frac{x_1^2}{\sum_{j=1}^n \Sigma_{jj}}.$$

این لم بیان می‌کند که توزیع توأم θ با عنوان تابعی از θ_i یا θ_{-i} ، ارزیابی شده در میانگین شرطی $E(\theta_{-i} | \theta_i)$ رفتار می‌کند. چون $\mathbf{x}$ گوسی نیست، این محاسبه یک تقریب خواهد بود. برای محور $k = 1, \dots, m$ این الگوریتم میانگین شرطی $E(\theta_{-i} | \theta_i)$ را با گوسی بودن θ ، طوری محاسبه می‌کند که در θ_j خطی و تنها به مد θ^* و کوواریانس Σ در راهبرد توری وابسته است. سپس از لم ۱.۳.۴ برای دستیابی به چگالی پسین کناری تقریب شده θ_j در هر مسیر از این محور استفاده می‌کنیم. برای هر مسیر از این محور باید فقط سه نقطه از این چگالی را که به وسیله‌ی لم ۱.۳.۴ تقریب زده شده‌اند، ارزیابی شوند به طوری که برای محاسبه مشتق دوم کافی است و انحراف معیارهای σ_k^- و σ_k^+ معادله به دست می‌دهند.

۴.۳.۴ تقریب چگالی پسین کناری عناصر میدان پنهان

پس از انتخاب نقاط انتگرال گیری، برای آن که بتوان با استفاده از دو تقریب $\tilde{\pi}(\theta_k|\mathbf{y})$ و $\tilde{\pi}(x_i|\theta_k, \mathbf{y})$ ، چگالی (۴.۴) را تقریب زد، باید چگالی پسین کناری x_i ها را محاسبه کرد. از این رو، در ادامه، با استفاده از سه روش تقریب گوسی، لاپلاس و لاپلاس ساده شده^{۱۶}، تقریب پسین $\pi(x_i|\theta_k, \mathbf{y})$ را محاسبه می کنیم و محاسن و معایب هر یک را بیان خواهیم کرد.

الف: تقریب گوسی

استفاده از $\tilde{\pi}_G(\mathbf{x}|\theta_k, \mathbf{y})$ و استخراج کناری آن ساده ترین راه محاسبه تقریب $\pi(x_i|\theta_k, \mathbf{y})$ است. با توجه به محاسبه $\tilde{\pi}(\theta|\mathbf{y})$ در مرحله انتخاب نقاط انتگرال گیری θ_k و همچنین محاسبه تقریب گوسی $\tilde{\pi}_G(\mathbf{x}|\theta_k, \mathbf{y})$ ، مقدار بردار میانگین نیز به دست آمده و می توان واریانس کناری را با استفاده از وارون ماتریس دقت Q به دست آورد. از آن جایی که معمولاً بعد میدان پنهان بزرگ است، وارون کردن ماتریسی با ابعاد بزرگ دشوار به نظر می آید. **رو و مارتینونو** (۲۰۰۷) رابطه بازگشتی را بر اساس معادلات تاکاهاشی^{۱۷} (تاکاهاشی و همکاران، ۱۹۷۳) معرفی کردند. این رابطه قادر است تا واریانس کناری را برای GMRF هایی با ابعاد بزرگ به سرعت محاسبه کند. بنابراین تقریب

$$\tilde{\pi}_G(x_i|\theta_k, \mathbf{y}) \sim N(\mu_i(\theta_k), \sigma_i^2(\theta_k)) \quad (۱۶.۴)$$

که در آن $\sigma_i^2(\theta_k)$ واریانس کناری مولفه x_i است، که برای $\pi(x_i|\theta_k, \mathbf{y})$ عموماً نتایج قابل قبولی را ارائه می دهد.

اگر $\mathbf{Q} = \mathbf{L}\mathbf{L}^T$ و $\mathbf{z} \sim N(\mathbf{o}, \mathbf{I})$ باشد، آن گاه $\mathbf{x}$ در معادله $\mathbf{L}^T \mathbf{x} = \mathbf{z}$ دارای توزیع $N(\mathbf{o}, \Sigma)$ است. لازم به ذکر است که می توان معادله $\mathbf{L}^T \mathbf{x} = \mathbf{z}$ را به صورت

$$l_{ii}x_i = z_i - \sum_{k=i+1}^n L_{ki}x_k \quad i = 1, \dots, n$$

هم نوشت. با ضرب طرفین این معادله در x_j به ازای $j \geq i$ و امید گرفتن از آن، درایه های ماتریس Σ توسط رابطه

$$\Sigma_{ij} = \frac{\delta_{ij}^2}{L_{ii}} - \frac{1}{L_{ii}} \sum_{k=i+1}^n L_{ki} \Sigma_{kj}, \quad j \geq i, i = n, \dots, 1 \quad (۱۷.۴)$$

محاسبه می شوند که در آن $\delta_{ii} = 1$ و اگر $i = j$ آن گاه $\delta_{ii} = 0$ خواهد بود. اما اگر $j \neq i$ ، معادله (۱۷.۴) دقیقاً معادل با معادله تاکاهاشی است که قادر به محاسبه سریع عناصر معینی از وارون یک ماتریس با استفاده از تجزیه چولسکی آن می باشد.

^{۱۶}Simplified Laplace approximation

^{۱۷}Takahashi equation

بدین ترتیب اگر $Q = LL^T = VDV^T$ باشد، به گونه‌ای که $L = VD^{\frac{1}{2}}$ در نتیجه D یک ماتریس قطری و V یک ماتریس پایین مثلثی با درایه‌های قطری یک است، بنابراین خواهیم داشت

$$\Sigma = D^{-1}V^{-1} + (I - V^T)\Sigma \quad (۱۸.۴)$$

چرا که $Q\Sigma = I$ بنابراین $VDV^T\Sigma = I$ پس

$$(VD)^{-1}VDV^T\Sigma = (VD)^{-1}.$$

بنابراین

$$IV^T\Sigma = D^{-1}V^{-1}.$$

پس

$$D^{-1}V^{-1} - V^T\Sigma = 0.$$

در نتیجه

$$D^{-1}V^{-1} - V^T\Sigma + \Sigma = \Sigma.$$

شکل ۵.۴ نواحی غیر صفر را برای رابطه (۱۸.۴) نشان می‌دهد. همان‌طور که ملاحظه می‌شود $D^{-1}V^{-1}$ یک ماتریس پایین مثلثی است، به‌طوری‌که $(D^{-1})_{ii} = (D^{-1}V^{-1})_{ii}$.

$$\Sigma = D^{-1}V^{-1} + \left[\begin{array}{c|c} I & \\ \hline & V^T \end{array} \right] \Sigma$$

$$= \left[\begin{array}{c|c} & \\ \hline & \end{array} \right] + \left[\begin{array}{c|c} & \\ \hline & \end{array} \right] \Sigma$$

شکل ۵.۴: شماتیک ماتریسی معادله‌ی تاکاهاشی

از این‌رو خوشبختانه نیازی به محاسبه‌ی V^{-1} نخواهیم داشت. بنابراین درایه‌های ماتریس Σ به‌صورت

$$\Sigma_{ij} = D_{ij}^{-1} - \sum_{k>j}^n V_{kj} \Sigma_{ik}, \quad i \geq j$$

محاسبه شده که پایه‌ای را برای محاسبه‌ی واریانس‌های کناری x_1 تا x_n فراهم می‌کند.

ب: تقریب لاپلاس

یکی از مشکلات به کارگیری تقریب گوسی، آن است که خطای ناشی از تشخیص نادرست مُد یا در نظر نگرفتن چولگی توزیع، ممکن است مقادیر آن را تحت تاثیر قرار دهد. برای رفع این مشکل، تقریب لاپلاس به صورت

$$\tilde{\pi}_{LA}(x_i|\theta, y) \propto \frac{\pi(\mathbf{x}, \theta, y)}{\tilde{\pi}_{GG}(\mathbf{x}_{-i}|x_i, \mathbf{x}, \theta, y)} \Big|_{\mathbf{x}_{-i}} = \mathbf{x}_{-i}^*(x_i, \theta) \quad (19.4)$$

معرفی شد که در آن $\tilde{\pi}_{GG}$ تقریب گوسی و $\mathbf{x}_{-i}^*(x_i, \theta)$ مد چگالی شرطی $\pi(\mathbf{x}_{-i}|x_i, \theta, y)$ است. همچنین لازم به ذکر است که $\tilde{\pi}_{GG}$ متفاوت از چگالی شرطی متناظر با $\tilde{\pi}(\mathbf{x}|\theta, y)$ است. در $\tilde{\pi}$ ماتریس دقت نسبت به x_i ثابت و میانگین آن تابعی خطی از x_i است. در حالی که در $\tilde{\pi}_{GG}$ ابتدا مُد $\mathbf{x}_{-i}^*(x_i, \theta)$ و سپس لگارتیم درست‌نمایی حول آن محاسبه می‌شود که ماتریس دقت آن با تغییر x_i تغییر می‌کند. از آن جا که چگالی $\tilde{\pi}_{GG}$ بر اساس شرطی کردن روی x_i و سپس به کار بردن تقریب لاپلاس بر روی سایر متغیرها قابل محاسبه است، از تقریب (۱۶.۴) که تنها بر اساس تقریب گوسی است، دقیق‌تر می‌باشد.

برای محاسبه‌ی عبارت (۱۹.۴) بایستی $\tilde{\pi}_{GG}$ را برای x_i و θ محاسبه کنیم که امری زمان‌بر و طاقت‌فرسا است. برای حل این مشکل، رو و همکاران (۲۰۰۹) دو تعدیل را برای تسهیل در محاسبات ارائه نمودند. در تعدیل اول به دلیل آن که میانگین شرطی $E[\mathbf{x}_{-i}|x_i, \theta, y]$ و مُد شرطی $\mathbf{x}_{-i}^*(x_i, \theta)$ به شرط گوسی بودن توزیع $(\mathbf{x}_{-i}^*|x_i, \theta, y)$ بر یکدیگر منطبق خواهند بود و به دلیل در نظر گرفتن توزیع پیشین گوسی برای $\pi(\mathbf{x}, \theta)$ ، گوسی بودن توزیع $(\mathbf{x}_{-i}^*|x_i, \theta, y)$ نیز چندان دور از انتظار نخواهد بود. بنابراین با استفاده از

$$\mathbf{x}_{-i}^*(x_i, \theta) \approx E\tilde{\pi}_G(\mathbf{x}_{-i}|x_i) \quad (20.4)$$

از میانگین به جای مُد استفاده و از گام بهینه‌سازی $\tilde{\pi}_{GG}(\mathbf{x}_{-i}|x_i, \theta, y)$ صرف نظر می‌کنیم. **هاسیو و همکاران (۲۰۰۴)** نشان دادند که استفاده از $E\tilde{\pi}(\mathbf{x}_{-i}|x_i)$ به جای $\mathbf{x}_{-i}^*(x_i, \theta)$ نتایج تقریباً مشابهی را در پی خواهد داشت.

در تعدیل دوم از این ایده استفاده می‌شود که تنها x_i هایی که در همسایگی x_i هستند، بر متغیر کناری x_i اثر می‌گذارند. اگر نزدیکی بین x_i و x_j به وابستگی بین دو رأس i و j تعبیر شود، تنها آن x_j هایی که در شعاع مشخصی از رأس i هستند در محاسبه‌ی چگالی کناری x_i لحاظ خواهند شد که این شعاع همسایگی با $R(\theta)$ نشان داده می‌شود. با استفاده از (۲۰.۴) رابطه

$$\frac{E\tilde{\pi}_G(x_j|x_i) - \mu_j(\theta)}{\sigma_j(\theta)} = a_{ij}(\theta) \frac{x_i - \mu_i(\theta)}{\sigma_j(\theta)} \quad (21.4)$$

برای بعضی a_{ij} که $i \neq j$ برقرار است. از این رو یک روش برای تعیین شعاع همسایگی، ساخت مجموعه‌ی $R(\theta)$ به صورت $R(\theta) = \{j : |a_{ij}| > 0.001\}$ است.

ج: تقریب لاپلاس ساده شده

در تقریب لاپلاس، حتی اگر از میانگین به جای مُد استفاده کنیم، همچنان فرآیند محاسباتی سنگین و زمان بر خواهد بود. با توجه به (۱۰.۴) و $\pi(\mathbf{x}_{-i}|x_i) = \frac{\pi(\mathbf{x})}{\pi(x_i)} \propto \pi(\mathbf{x})$ مخرج کسر (۱۹.۴) به صورت

$$\begin{aligned} \tilde{\pi}_{GG}(\mathbf{x}_{-i}|x_i, \boldsymbol{\theta}, \mathbf{y}) \Big|_{x_i=E_{\tilde{\pi}_G}(\mathbf{x}_{-i}|x_i)} &\propto (\gamma\pi)^{-\frac{n}{\gamma}} |\mathbf{Q} + \text{diag}(\mathbf{c})|^{\frac{1}{\gamma}} \\ &\times \exp\left(-\frac{1}{\gamma} \mathbf{x}^T \mathbf{Q} \mathbf{x} + b^T \mathbf{x} - \frac{1}{\gamma} \mathbf{x}^T (\text{diag}(\mathbf{c}) \mathbf{x})\right) \Big|_{\mathbf{x}_{-i}=E_{\tilde{\pi}_G}(\mathbf{x}_{-i}|x_i)} \end{aligned} \quad (22.4)$$

است، که در آن $c_i = -\frac{\partial^2 \pi(y_j|x_j, \boldsymbol{\theta})}{\partial x_j^2}$. بنابراین برای هر $i = 1, \dots, n$ ماتریس دقت $\mathbf{x}_{-i}|x_i$ به صورت زیرماتریس اصلی $Q_{[-i,-j]}$ است. از طرفی رابطه

$$\begin{aligned} \pi(\mathbf{x}) &= \pi(x_i) \pi(\mathbf{x}_{-i}|x_i) \\ &\propto \text{Var}(x_i)^{-\frac{1}{\gamma}} |Q_{[-i,-j]}|^{\frac{1}{\gamma}} \exp\left(-\frac{1}{\gamma} \mathbf{x}^T \mathbf{Q} \mathbf{x}\right) \end{aligned}$$

برقرار است. بنابراین رابطه‌ی (۲۲.۴) را می‌توان به صورت

$$\begin{aligned} \tilde{\pi}_{GG}(\mathbf{x}_{-i}|x_i, \boldsymbol{\theta}, \mathbf{y}) \Big|_{\mathbf{x}_{-i}=E_{\tilde{\pi}_G}(\mathbf{x}_{-i}|x_i)} &\propto (\gamma\pi)^{-\frac{n}{\gamma}} \text{Var}(x_i)^{-\frac{1}{\gamma}} |Q_{[-i,-i]} + \text{diag}(\mathbf{c})|^{\frac{1}{\gamma}} \exp\left(-\frac{1}{\gamma} (x_i, E(\mathbf{x}_{-i}|x_i)^T) Q \begin{pmatrix} x_i \\ E(\mathbf{x}_{-i}|x_i) \end{pmatrix}\right) \\ &\exp\left(b^T \begin{pmatrix} x_i \\ E(\mathbf{x}_{-i}|x_i) \end{pmatrix} - \frac{1}{\gamma} (x_i, E(\mathbf{x}_{-i}|x_i)^T) (\text{diag}(\mathbf{c})) \begin{pmatrix} x_i \\ E(\mathbf{x}_{-i}|x_i) \end{pmatrix}\right) \\ &\propto |Q_{[-i,-i]} + \text{diag}(\mathbf{c})|^{\frac{1}{\gamma}} \exp\left(-\frac{1}{\gamma} (x_i, E(\mathbf{x}_{-i}|x_i)^T) (\text{diag}(\mathbf{c})) \begin{pmatrix} x_i \\ E(\mathbf{x}_{-i}|x_i) \end{pmatrix}\right) \end{aligned}$$

نوشت. **رو و همکاران (۲۰۰۹)**، نشان دادند که اگر $\mathbf{x} \sim N(\circ, \Sigma)$ ، آن گاه

$$-\frac{1}{\gamma} (x_i, E(\mathbf{x}_{-i}|x_i)^T) \Sigma^{-1} \begin{pmatrix} x_i \\ E(\mathbf{x}_{-i}|x_i) \end{pmatrix} = -\frac{1}{\gamma} \frac{x_i^2}{\Sigma_{ii}}, \quad i = 1, \dots, n.$$

بنابراین

$$\log \tilde{\pi}_{GG}(\mathbf{x}_{-i}|x_i, \boldsymbol{\theta}, \mathbf{y}) \Big|_{\mathbf{x}_{-i}=E_{\tilde{\pi}_G}(\mathbf{x}_{-i}|x_i)} \propto c_0 + \frac{1}{\gamma} \log |Q^* + \text{diag}(\mathbf{c})|. \quad (23.4)$$

در محاسبه‌ی دترمینان موجود در (۲۳.۴) باید برای هر x_i یک ماتریس $(N-1) \times (N-1)$ محاسبه شود که امری بسیار زمان بر است. **رو و همکاران (۲۰۰۹)** برای حل این مشکل تقریب لاپلاس ساده شده را معرفی کردند. در این روش با تعریف $x_i^{(s)} = \frac{x_i - \mu_i(\boldsymbol{\theta})}{\sigma_i(\boldsymbol{\theta})}$ صورت کسر (۱۹.۴) حول $x_i = \mu_i(\boldsymbol{\theta})$ به صورت

$$\begin{aligned} \log\{\pi(\mathbf{x}, \boldsymbol{\theta}, \mathbf{y})\} \Big|_{\mathbf{x}_{-i}=E_{\tilde{\pi}_G}(\mathbf{x}_{-i}|x_i)} &= \\ &-\frac{1}{\gamma} (x_i^{(s)})^2 + \frac{1}{\gamma} (x_i^{(s)})^3 \sum_{j \in \frac{\mathbb{Z}}{i}} d_j^{(\gamma)} \{\mu_i(\boldsymbol{\theta}), \boldsymbol{\theta}\} \{\sigma_i(\boldsymbol{\theta}) a_{ij}(\boldsymbol{\theta})\}^3 + \dots \end{aligned} \quad (24.4)$$

بسط داده می‌شود. جملات مرتبه‌ی اول و دوم تقریب گوسی و جمله‌ی مرتبه‌ی سوم تصحیحی برای چولگی فراهم می‌کند. در بسط لگارتیم مخرج کسر (۱۹.۴) با توجه به این که برای هر ماتریس M رابطه‌ی $\frac{\partial \log |M|}{\partial m} = \text{Trace}(M^{-1} \frac{\partial M}{\partial m})$ برقرار است (ماردیا و همکاران، ۱۹۸۹)، رابطه‌ی

$$\begin{aligned} \frac{d \log |Q^* + \text{diag}(\mathbf{c})|}{dx_i} &= \text{Trace}[Q^* + \text{diag}(\mathbf{c})]^{-1} \frac{d[Q^* + \text{diag}(\mathbf{c})]}{dx_i} \\ &= \text{Trace}[Q^* + \text{diag}(\mathbf{c})]^{-1} \text{diag}[d^{\text{r}}(x_i, \theta)] \\ &= \sum_j \text{Var}(x_j | x_i) d_j^{\text{r}}(x_i, \theta) \\ &= \sum_j \sigma_j(\theta) \{1 - \text{corr}_{\tilde{\pi}_G}(x_i, x_j | \theta)^2\} d_j^{\text{r}}(x_i, \theta) \end{aligned} \quad (25.4)$$

به دست می‌آید، که در آن $\text{Var}(x_j | x_i) = \sigma_i(\theta) \{1 - \text{corr}_{\tilde{\pi}_G}(x_i, x_j | \theta)^2\}$. با استفاده از رابطه‌ی (۲۵.۴)، بسط تیلور مخرج کسر (۱۹.۴) حول $x_i = \mu_i(\theta)$ به صورت

$$\begin{aligned} \log \{ \pi(\mathbf{x}, \theta, \mathbf{y}) \} |_{\mathbf{x}_{-i} = E_{\tilde{\pi}_G}(\mathbf{x}_{-i} | x_i)} = \\ c_0 - \frac{1}{\varphi} (X_i^{(s)}) \sum_{j \in \frac{\mathcal{I}}{i}} \sigma_j(\theta) \{1 - \text{corr}_{\tilde{\pi}_G}(x_i, x_j | \theta)^2\} d_j^{\text{r}}(\mu_i(\theta)_i, \theta) \sigma_j(\theta) a_{ij}(\theta) + \dots \end{aligned} \quad (26.4)$$

حاصل می‌شود که در آن c_0 مقداری ثابت است. با توجه به (۲۴.۴) و (۲۶.۴) خواهیم داشت

$$\gamma^{(1)}(\theta) = \frac{1}{\varphi} (x_i^{(s)}) \sum_{j \in \frac{\mathcal{I}}{i}} \sigma_i(\theta) \{1 - \text{corr}_{\tilde{\pi}_G}(x_i, x_j | \theta)^2\} d_j^{\text{r}}(\mu_i(\theta)_i, \theta) a_{ij}(\theta)$$

و

$$\gamma^{(3)}(\theta) = \sum_{j \in \frac{\mathcal{I}}{i}} d_j^{\text{r}}(\mu_i(\theta), \theta) \{ \sigma_j(\theta) a_{ij}(\theta) \}^3.$$

بنابراین تقریب لاپلاس ساده شده به صورت زیر به دست می‌آید:

$$\tilde{\pi}_{SLA} = c_0 - \frac{1}{\varphi} (x_i^{(s)}) + \gamma_i^{(1)}(\theta) (x_i^{(s)}) + \frac{1}{\varphi} + (x_i^{(s)}) \gamma_i^{(3)}(\theta) + \dots$$

۴.۴ ملاک‌های ارزیابی

یکی از راه‌های متداول برای مقایسه مدل‌ها در تحلیل بیزی، استفاده از عامل بیزی^{۱۸} است. فرض کنید بردار مشاهدات $\mathbf{y}$ را داریم و هدف ما مقایسه دو مدل

$$M_1 : \pi_1(\mathbf{y} | \theta_1), \quad M_2 : \pi_2(\mathbf{y} | \theta_2)$$

^{۱۸}Bayes factor

باشد. توزیع های پیشین $\pi_1(\theta_1)$ و $\pi_2(\theta_2)$ را برای پارامترها و احتمالات پیشین $\pi(M_1)$ و $\pi(M_2)$ را برای مدل ها در نظر می گیریم. بر اساس قانون بیز، بخت پسین ها برای دو مدل به صورت

$$\frac{\pi(M_1|\mathbf{y})}{\pi(M_2|\mathbf{y})} = \frac{\pi(M_1)}{\pi(M_2)} \cdot \frac{\int_{\theta_1} \pi_1(\mathbf{y}|\theta_1)\pi_1(\theta_1)d\theta_1}{\int_{\theta_2} \pi_2(\mathbf{y}|\theta_2)\pi_2(\theta_2)d\theta_2}$$

عامل بیزی $\times$ بخت پیشین مدل ها =

است. بعد از مرتب سازی، عامل بیزی به صورت

$$\begin{aligned}\beta(\mathbf{y}) &= \frac{\pi(M_1|\mathbf{y})}{\pi(M_2|\mathbf{y})} \cdot \frac{\pi(M_2)}{\pi(M_1)} \\ &= \frac{\pi(M_1|\mathbf{y})/\pi(M_2|\mathbf{y})}{\pi(M_1)/\pi(M_2)} \\ &= \frac{\text{بخت پسین مدل ها}}{\text{بخت پیشین مدل ها}}\end{aligned}$$

به دست می آید. اگر $\pi(M_1) = \pi(M_2)$ و فضای پارامتر θ_1 و θ_2 یکسان باشند، آن گاه

$$\begin{aligned}\beta(\mathbf{y}) &= \frac{\pi(M_1|\mathbf{y})}{\pi(M_2|\mathbf{y})} \cdot \frac{\pi(M_2)}{\pi(M_1)} \\ &= \frac{\frac{\pi(M_1, \mathbf{y})}{\pi(\mathbf{y})\pi(M_1)}}{\frac{\pi(M_2, \mathbf{y})}{\pi(\mathbf{y})\pi(M_2)}} \\ &= \frac{\pi(\mathbf{y}|M_1)}{\pi(\mathbf{y}|M_2)}.\end{aligned}$$

یعنی عامل بیزی با نسبت درست نمایی ها برابر می شود. بنابراین مدل ها را می توان با استفاده از درست نمایی های دو مدل، یعنی $\pi(\mathbf{y}|M_1)$ و $\pi(\mathbf{y}|M_k)$ مقایسه کرد. برای این منظور تقریب درست نمایی های $\pi(\mathbf{y}|M_k)$ با استفاده از (۹.۴)، به روش INLA به صورت

$$\hat{\pi}(\mathbf{y}|M_k) = \int \frac{\tilde{\pi}_k(\mathbf{x}, \boldsymbol{\theta}, \mathbf{y})}{\tilde{\pi}_G(\mathbf{x}|\boldsymbol{\theta}, \mathbf{y})} d\boldsymbol{\theta}, \quad k = 1, 2 \quad (27.4)$$

محاسبه می شود.

۱.۴.۴ ملاک اطلاع انحراف

برای مقایسه مدل های سلسله مراتبی پیچیده از ملاک اطلاع انحراف^{۱۹} (DIC) که اندازه ای از پیچیدگی^{۲۰} (تعداد موثر پارامترها) و برازش^{۲۱} مدل است، استفاده می شود. با توجه به این که افزایش تعداد پارامترها تأثیر مستقیمی بر پیچیدگی و قدرت پیشگویی مدل دارد، بهتر است به هنگام مقایسه مدل ها تعداد پارامترهای موثر نیز مورد توجه قرار گیرد. **دمپستر (۱۹۷۴)**،

^{۱۹}Deviance information criterion

^{۲۰}Complexity

^{۲۱}Fit

ملاکی تحت عنوان انحراف را بر اساس لگاریتم درستنمایی توزیع پسین معرفی کرد که از آن برای استخراج کمیت‌های برازش و پیچیدگی استفاده می‌شود و نهایتاً با ترکیبی از این دو ملاک DIC تعریف می‌شود (اشپیگل هالترو و همکاران، ۲۰۰۲).

فرض کنید بردار مشاهدات y در اختیار باشد و هدف انتخاب بهترین مدل از بین مدل‌های برازش‌یافته به این مجموعه باشد. با استفاده از رهیافت بیزی، توزیع توأم $\pi(y, \theta)$ به صورت

$$\pi(y, \theta) = \pi(y|x)\pi(x|\theta)\pi(\theta)$$

است که در آن $\phi = (x, \theta)$ و با شرط معلوم بودن x و θ ، y_i ‌ها مستقل شرطی هستند. **دمیستر** (۱۹۷۴)، ملاک انحراف را به صورت

$$D(x) = -2 \log \pi(y|x) + 2 \log \pi(y) \quad (28.4)$$

تعریف کرد، که در آن $\pi(y)$ ، عبارت استاندارد است که فقط تابعی از داده‌ها است و بنابراین تأثیری در مقایسه مدل‌ها ندارد. لازم به ذکر است که عبارت (۲۸.۴) را انحراف بیزی^{۲۲} نیز می‌نامند. **اشپیگل هالترو و همکاران** (۲۰۰۲)، کمیت برازش را بر اساس میانگین پسین انحرافات به صورت

$$\bar{D} = E_{x|y}[D]$$

و کمیت پیچیدگی را اختلاف بین میانگین پسین انحرافات و انحراف بر اساس میانگین پسین پارامترها به صورت

$$\begin{aligned} pD &= E_{x|y}[D] - D[E_{x|y}(x)] \\ &= \bar{D} - D(\bar{x}) \end{aligned}$$

تعریف کردند. ملاک DIC با در نظر گرفتن دو مفهوم برازش و پیچیدگی به صورت

$$\begin{aligned} DIC &= \bar{D} + pD \\ &= D(\bar{x}) + 2pD \end{aligned}$$

تعریف می‌شود. مدل‌های با DIC کوچکتر، داده‌ها را بهتر پشتیبانی می‌کنند.

۲.۴.۴ ملاک پیش‌گویی شرطی مولفه‌ها

برای مقایسه مدل‌ها، از نظر اعتبار پیش‌گویی می‌توان از ملاک پیش‌گویی شرطی مولفه‌ها^{۲۳} (CPO) برای مشاهدات به صورت

$$CPO_i = \pi(y_i | y_{-i}), \quad i = 1, \dots, n$$

^{۲۲} Bayesian deviance

^{۲۳} Conditional predictive ordinates

استفاده کرد، که در آن y_{-i} همان بردار y به‌جز درایه i ام است. مقادیر صفر و خیلی کوچک CPO بیان‌گر پرت بودن y_i متناظر است. از منهای میانگین لگارتیم CPO_i ها برای مقایسه مدل‌ها استفاده می‌شود، که ملاک نمره لگاریتمی اعتبارسنجی متقابل^{۲۴} (LogScore) نامیده می‌شود (گنتینگ و رفتی، ۲۰۰۷). با توجه به این که مقادیر بزرگ CPO_i ها گویای پیش‌گویی برتر مدل است، بنابراین هرچه مقدار LogScore کوچک باشد، مدل برای پیش‌گویی مناسب‌تر خواهد بود.

۵.۴ داده‌های فضایی سه بعدی

میدان‌های تصادفی گوسی (معروف به GRF) از مولفه‌های مهم مدل‌بندی در آمار فضایی هستند و نقش برجسته‌ای به‌ویژه در زمین‌آمار دارند (کرس، ۱۹۹۳؛ استین، ۱۹۹۹؛ دیگل و ریبریو، ۲۰۰۷؛ چلیس و دلفینر، ۲۰۱۲). همچنین GRFها در مدل‌های فضایی سلسله‌مراتبی مدرن، یک مولفه ساختاری مهم را تشکیل می‌دهند (بارتلت، ۱۹۳۸) که ویژگی‌های اصلی میدان فضایی توسط تابع کوواریانس آن‌ها بیان می‌شود. میدان‌های تصادفی گوسی از محدود مدل‌های چند متغیره‌ی مناسب با قابلیت محاسبه‌ی ثابت نرمال‌ساز هستند که از ویژگی‌های تحلیلی خوبی نیز برخوردارند. اما از منظر محاسباتی، GRFها برای حجم‌های نمونه بزرگ با مشکل مواجه می‌شوند، زیرا هزینه‌ی تجزیه ماتریس‌های چگال (غیرتُنک) با بعد بالا می‌تواند بسیار گزاف باشد. امروزه اگرچه قدرت محاسباتی در بالاترین سطح ممکن خود قرار دارد، اما در واقعیت هنوز کندی محاسبات سد راه بسیاری از پژوهش‌های کاربردی است.

قدرت محاسباتی و کارایی روش INLA زمانی است که میدان تصادفی فضایی، مارکوفی باشد و بتوان ماتریس دقت (که عکس ماتریس کوواریانس میدان است) را تُنک در نظر گرفت. بنابراین، این روش برای داده‌های زمین‌آمار به‌ویژه زمانی که حجم داده‌ها بزرگ است، نمی‌تواند کارا باشد. برای رفع این مشکل و استفاده از روش تقریبی INLA برای تحلیل داده‌های زمین‌آمار، لیندگرین و همکاران (۲۰۱۱) روشی را با عنوان معادلات دیفرانسیل جزئی تصادفی^{۲۵} (SPDE) معرفی کردند که با مکانیسمی یک GRF را با یک تقریب GMRF به‌صورت تقریبی بازنویسی می‌کند. با روش پیشنهادی آن‌ها می‌توان یک مدل GRF را برای داده‌های زمین‌آمار توسط یک مدل مارکوفی GMRF نمایش داد و سپس از روش INLA برای برازش آن به داده‌ها بهره برد. به این شکل، کارایی محاسباتی روش تقریبی بیزی با ترکیب رهیافت SPDE حفظ می‌شود.

پس، هدف اصلی رهیافت SPDE یافتن یک فرآیند تصادفی مارکوفی GMRF، با همسایگی فضایی و ماتریس دقت تُنک Q توسط گسسته‌سازی ناحیه فضایی پیوسته است. با توجه به تمرکز پایان‌نامه بر تحلیل داده‌های زمین‌آمار، میدان نفتی، از رهیافت ترکیبی INLA+SPDE

^{۲۴}Cross Validated Logarithmic Score

^{۲۵}Stochastic partial differential equation

برای مدل‌بندی آن‌ها استفاده می‌کنیم. استفاده از این رهیافت برای داده‌های فضایی سه بعدی به ندرت و در چارچوب‌های غیرمنسجم با ابزارهای غیرمتحد در نرم‌افزارهای مختلف انجام شده است. به‌عنوان نمونه می‌توان به **ژانگ و همکاران (۲۰۱۶)** اشاره کرد. ما در این پایان‌نامه، برای اولین بار از رهیافت ترکیبی مذکور و مختص داده‌های سه بعدی برای تحلیل داده‌های فضایی استفاده می‌کنیم. برای این کار، ابتدا رهیافت SPDE را معرفی می‌کنیم.

۶.۴ رهیافت معادلات دیفرانسیل جزئی تصادفی

فرض کنید $\{\omega(s), s \in D \subseteq \mathbb{R}^d\}$ یک میدان تصادفی گوسی مانای مرتبه دوم با تابع کوواریانس مترن (معرفی‌شده در رابطه (۸.۱) در فصل اول) با پارامتر همواری k و مقیاس ϕ است. افزون بر این، فرض کنید فرآیند $\omega(s)$ در موقعیت‌های $s_i, i = 1, \dots, n$ ، مشاهده شده است. روش SPDE یک GMRF با همسایگی فضایی و ماتریس Q می‌یابد که در بهترین حالت نشان‌دهنده یک GRF مترن است. در این رهیافت، میدان تصادفی مترن به‌صورت یک ترکیب خطی از توابع تکه‌ای خطی که بر یک مثلث‌بندی از دامنه‌ی D تعریف شده است، بیان می‌شود. به‌عبارت ساده‌تر دامنه فضای D به مجموعه‌ای از مثلث‌ها تقسیم‌بندی می‌شود که حداکثر در یک یال یا زاویه‌ی مشترک متقاطع اشتراک دارند. موقعیت‌های $s_1, \dots, s_n$ رؤس مثلث‌بندی اولیه در نظر گرفته می‌شوند و سپس دیگر رؤس به‌منظور دست یافتن به یک مثلث‌بندی مناسب برای پیش‌گویی فضایی و سایر اهداف اضافه می‌شوند.

لیندگرین و همکاران (۲۰۱۱) نشان دادند تنها جواب‌های پایدار برای یک SPDE میدان‌های تصادفی مترن هستند، به‌طوری که یک میدان تصادفی گوسی $\omega(s)$ با تابع کوواریانس مترن و پارامترهای هموارسازی k و مقیاس ϕ تعریف می‌شود، که در آن $\sigma^2 = \frac{\Gamma(k)}{\Gamma(k+\frac{d}{\phi})(\frac{d}{\phi}\pi)^{\frac{d}{\phi}}\pi^{\frac{d}{\phi}}}$ و $K_k(\cdot)$ تابع بسل اصلاح شده نوع دوم است (**ویتل ۱۹۵۴ و ۱۹۵۴**). یک جواب پایدار برای معادله‌ی دیفرانسیل جزئی تصادفی به‌صورت

$$(\phi^2 - \Delta^2)^{\frac{\alpha}{2}} \tau(s) \omega(s) = W(s), \quad s \in \mathbb{R}^d, \alpha = k + \frac{d}{\phi}, \phi > 0, k > 0 \quad (29.4)$$

است، که در آن W میدان تصادفی نوفه سفید^{۲۶}، Δ عمل‌گر دیفرانسیلی مرتبه دوم یا همان عمل‌گر لاپلاس^{۲۷}، $\Delta = \sum_{i=1}^d \sigma^2 / \omega_i^2$ ، و τ پارامتر کنترل‌گر واریانس است. عبارت $(\phi^2 - \Delta^2)^{\frac{\alpha}{2}}$ جمله‌ی عمل‌گر دیفرانسیل^{۲۸} است. با به کار بردن تعریف فوریه، یک لاپلاس کسری^{۲۹} به شکل

$$F(\phi^2 - \Delta^2)^{\frac{\alpha}{2}}(\psi) = (\phi^2 + \|K\|^2)^{\frac{\alpha}{2}}(F\psi)(K)$$

^{۲۶}White noise

^{۲۷}Laplacian operator

^{۲۸}Pesudo-differential operator

^{۲۹}Fractional Laplacian

در نظر گرفته می‌شود، که در آن ψ تابعی در $\mathbb{R}^d$ است. **لیندگرین و همکاران (۲۰۱۱)** ابتدا یک صورت ضعیف تصادفی^{۳۰} معادله‌ی دیفرانسیل جزئی (۲۹.۴) را برای یک میدان تصادفی مارکوفی گوسی که بیان‌گر میدان تصادفی مترن روی شبکه مثلث‌بندی باشد، در نظر گرفتند. شبکه‌بندی مثلثی نسبت به روش‌های دیگر مانند روش‌های مبتنی بر شبکه‌های معمول و سنتی، منعطف‌تر است.

یک روش عددی برای حل تقریبی معادلات دیفرانسیل جزئی، روش عناصر متناهی^{۳۱} است و مثلث‌بندی روشی معمول برای حل معادلات با استفاده از روش عناصر متناهی است (برنر و اسکات، ۲۰۰۷). در این روش دقت جواب‌ها به نحوه‌ی مثلث‌بندی ناحیه فضایی بستگی دارد. **لیندگرین و همکاران (۲۰۱۱)** برای نمایش عناصر متناهی SPDE، یک ترکیب خطی به صورت

$$\omega(\mathbf{s}) = \sum_{g=1}^G \psi_g(\mathbf{s}) \hat{\omega}_g \quad (30.4)$$

تعریف کردند، که در آن $\{\psi_g\}$ مجموعه‌ای از توابع پایه، $\hat{\omega}$ ها وزن‌های ترکیب خطی دارای توزیع گوسی با میانگین صفر و G تعداد کل رئوس مثلث‌بندی است. توابع پایه طوری انتخاب می‌شوند که در رأس g برابر یک و در سایر رئوس صفر باشند.

در بسته‌ی نرم‌افزاری R-INLA برای $d = 2$ معادله‌ی دیفرانسیل جزئی تصادفی (۲۹.۴) به‌صورت عمومی

$$(\phi^2(\mathbf{s}) - \Delta^2)^{\frac{\alpha}{2}} \tau(\mathbf{s}) \omega(\mathbf{s}) = W(\mathbf{s}), \quad \mathbf{s} \in \mathbb{R}^2, \alpha = k + \frac{d}{4}, \phi > 0, k > 0$$

در نظر گرفته می‌شود که میدان‌های تصادفی نامانا را نیز شامل می‌شود، یعنی میدان‌هایی که واریانس $W(\mathbf{s})$ و پارامتر ϕ در آن‌ها ثابت نیست و به موقعیت فضایی وابسته است. البته برای سادگی، واریانس $W(\mathbf{s})$ را ثابت و پارامتر مقیاس را برای $\omega(\mathbf{s})$ به شکل $\tau(\mathbf{s})$ در نظر می‌گیرند. در این صورت، اگر یک یا هر دو پارامتر τ و ψ ثابت نباشند، میدان نامانا خواهد شد و بر اساس رابطه‌ی (۸.۱) واریانس کناری میدان به شکل $\sigma^2 = \frac{1}{4\pi\tau\phi^2}$ تغییر می‌یابد. **لیندگرین و همکاران (۲۰۱۱)** نشان دادند وزن‌های گوسی در نظر گرفته‌شده در (۳۰.۴) یک ساختار مارکوفی دارند و میدان تصادفی مترن را به یک GMRF تبدیل می‌کند. ماتریس دقت GMRF تولیدشده، وقتی $\alpha = 2$ باشد، برای وزن‌های $\omega = \{\hat{\omega}_1, \dots, \hat{\omega}_G\}$ با استفاده از شرایط مرزی نیومن (چنگ و چنگ، ۲۰۰۵) به‌صورت

$$\mathbf{Q} = \tau(\phi^2 \mathbf{C} + 2\phi^2 \mathbf{G} + \mathbf{G} \mathbf{C}^{-1} \mathbf{G}) \quad (31.4)$$

است، که در آن درایه‌های ماتریس $\mathbf{C}$ به‌صورت

$$C_{ii} = \int \psi_i(\mathbf{s}) d\mathbf{s}$$

^{۳۰} Stochastic weak formulation

^{۳۱} Finite element method

و درایه‌های ماتریس تنگ G به صورت

$$\mathbf{G}_{ij} = \int \nabla \psi_i(\mathbf{s}) \nabla \psi_g(\mathbf{s}) d\mathbf{s}$$

است، که در آن منظور از ∇ عمل گرادیان است. ماتریس دقت $\mathbf{Q}$ تُنک است و درایه‌های آن به پارامترهای ϕ و τ وابسته است. به این شکل، ω یک فرآیند تصادفی مارکوفی با توزیع نرمال با میانگین صفر و واریانس $\mathbf{Q}^{-1}$ است که برای اهداف محاسباتی کارای INLA مناسب خواهد بود.

۷.۴ رهیافت INLA+SPDE برای مدل‌بندی فضایی سه بعدی

با بازنویسی پیشگوی خطی (۱.۴) برای هر کدام از اثرات ثابت، غیرخطی ناپارامتری و اثرات تصادفی (شامل اثرات فضایی)، می‌توان پیشگوی خطی را به‌عنوان یک ترکیب خطی از مولفه‌های اثرات پنهان x نوشت. این بازنویسی کمک می‌کند تا بتوان در تقریب مبتنی بر SPDE برازش مدل با روش INLA را به‌صورتی واحد و کارا اجرا کرد.

به‌طور دقیق‌تر می‌توان پیشگوی (۱.۴) را به‌صورت زیر بازنویسی کرد:

$$\eta_i = \sum_k h_k(z_i^k)$$

که در آن k بیان گر k امین مولفه از پیشگو است و z_i^k مقدار متغیر تبیینی با اثر ثابت یا اثر (ناپارامتری) غیرخطی است یا متغیری است که اندیس مربوط به اثرات تصادفی را در خود جای داده است. تابع $h_k(\cdot)$ نیز تابعی است که z_i^k را به k امین مولفه این معادله پیوند می‌زند. در برخی موارد هر مشاهده y_i به ترکیب خطی مولفه‌های x وابسته است و توزیع مشاهدات پاسخ را می‌توان به صورت زیر بازنویسی کرد:

$$y_i | \mathbf{x}, \boldsymbol{\theta} \sim \pi \left(y_i | \sum_j \mathbf{A}_{ij} \mathbf{x}, \boldsymbol{\theta} \right).$$

مولفه‌های A_{ij} عناصر ماتریس A هستند که ماتریس مشاهده نامیده می‌شود. ماتریس مشاهده نگاشته بین میدان پنهان فضایی (تعریف شده روی شبکه مثلث‌بندی شده) و مشاهدات (تعریف شده در مجموعه‌ای از مکان‌ها) ایجاد می‌کند. روش R-INLA، پیش‌گوی خطی جدید η^* را به صورت ترکیب خطی η تولید می‌کند، یعنی

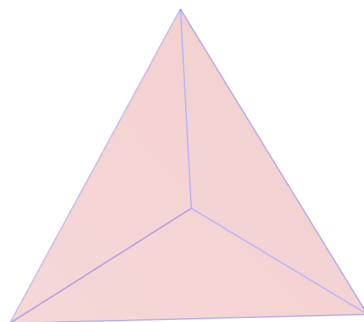
$$\eta^* = A\eta$$

و درست‌نمایی به‌وسیله‌ی η_i^* به‌جای η_i به میدان پنهان مرتبط می‌شود. بنابراین

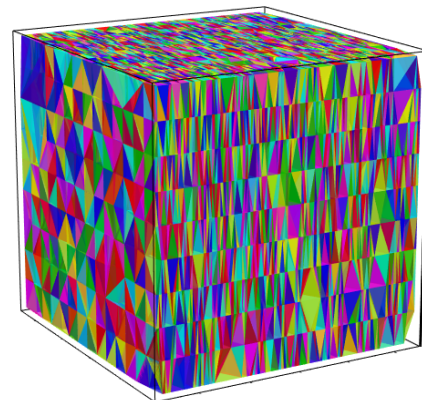
$$\pi(\mathbf{y}|\mathbf{x}, \boldsymbol{\theta}) = \prod_{i=1}^n \pi(y_i|\eta_i^*, \boldsymbol{\theta}).$$

۱.۷.۴ مثلث‌بندی ناحیه سه بعدی

برای اجرای روش SPDE ابتدا باید ناحیه تحت مطالعه را مثلث‌بندی کرد. در بسته نرم‌افزاری R-INLA برای حالت یک و دو بعدی، توابع تعریف‌شده‌ای برای این کار وجود دارند. اما در حالت عمومی سه بعدی تا کنون چنین ابزاری وجود نداشته است. با طرح مساله مورد نظر این پایان‌نامه با توسعه‌دهنده اصلی رهیافت SPDE، یعنی فین لیندگرین، یک بسته نرم‌افزاری جدید که به بسته INLA افزوده می‌شود، نوشته شد. در این بسته با نام `inlamesh3d` می‌توان مثلث‌بندی سه بعدی را که در آن از چهاروجهی^{۳۲} به‌جای مثلث استفاده می‌شود، اجرا کرد. برای این کار از روش مثلث‌بندی دلاونی^{۳۳} که در بسته نرم‌افزاری `geometry` موجود است، استفاده می‌شود. تصویر یک چهاروجهی منتخب از یک ناحیه مثلث‌بندی شده سه بعدی در ۶.۴ (آ و ب) نمایش داده شده است.



(ب)



(آ)

شکل ۶.۴: گسسته‌سازی ناحیه مثال شبیه‌سازی در فضای سه بعدی با چهاروجهی‌ها (آ) و یک چهاروجهی منتخب (ب)

ماتریس‌های ساختاری لازم برای روش عناصر متناهی، در این حالت، مشابه حالت دوبعدی قابل محاسبه هستند. این ماتریس‌های ساختاری براساس رابطه (۳۱.۴) برای ساخت ماتریس دقت میدان تقریبی لازم هستند. این کار را می‌توان به صورت سراسر انجام داد و محاسبات پیچیده‌ای را نیاز ندارد.

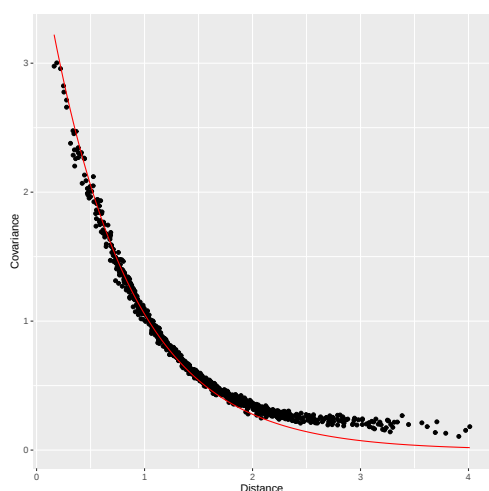
^{۳۲} Tetrahedron

^{۳۳} Delaunay triangulation

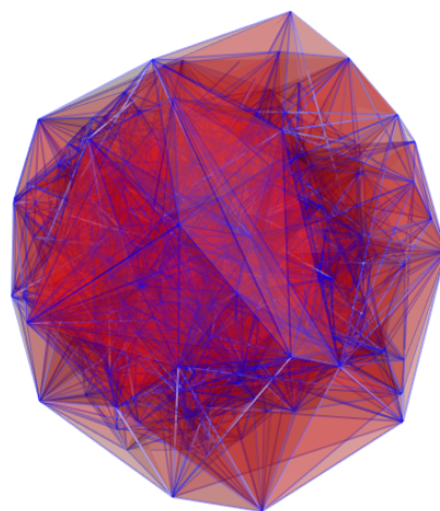
تعیین پارامترهای میدان مترن در اجرای روش SPDE نیز شبیه حالت دوبعدی و با استفاده از توابع استاندارد موجود در روش مذکور در بسته R-INLA قابل انجام است. برای ارزیابی روش پیشنهادی جدید از یک مثال شبیه‌سازی استفاده می‌کنیم.

مثال شبیه‌سازی

در این مثال نقاط مکانی سه بعدی به حجم 1000 را از توزیع نرمال سه متغیره استاندارد تولید کردیم. بنابراین در یک ناحیه سه بعدی، هزار نقطه مکانی نامنظم قرار می‌گیرند. مقادیر واقعی پارامترهای دامنه و واریانس را به ترتیب $1/5$ و 2 انتخاب کردیم. همچنین مقادیر 1 و 1 را به‌عنوان میانه توزیع‌های پیشین پارامترهای دامنه و واریانس میدان مترن در نظر گرفتیم. بر اساس این نقاط و ناحیه مثلث‌بندی شده (شکل ۷.۴ (آ) را ببینید)، ماتریس دقت برآورد شد و تابع کوواریانس میدان برآوردشده بر حسب فاصله نقاط محاسبه و در قسمت ب شکل ۷.۴ رسم شد. در این شکل منحنی تابع کوواریانس واقعی که یک تابع نمایی است نیز رسم شده است. همانطور که مشاهده می‌شود، به‌جز در مرزها (با کمی انحراف از واقعیت به دلیل محدودیت ناحیه) کوواریانس برآوردشده با واقعیت تطابق دارد. این نتیجه نشانی از درستی تقریب GMRF حاصل از گسسته‌سازی سه بعدی در ناحیه تحت مطالعه است.



(ب)



(آ)

شکل ۷.۴: (آ) گسسته‌سازی سه بعدی ناحیه شبیه‌سازی شده، (ب) منحنی تابع کوواریانس واقعی در مثال شبیه‌سازی به همراه تقریب حاصل از روش SPDE

برای ارزیابی توانایی بازیابی مقادیر واقعی پارامترهای میدان تصادفی تولیدشده، از میدان مترن با ماتریس دقت مذکور و میانگین Ax در 100 نقطه مکانی شبیه‌سازی شده، مقادیر پاسخ را تولید کردیم. به داده‌های تولید شده یک نوفه نرمال با انحراف معیار 0.1 نیز اضافه کردیم. نتایج برآوردهای میانه پسینی به همراه فواصل باورمند بیزی ۹۵ درصد در جدول ۴.۴ گزارش

شده‌اند. دو پارامتر تابع کوواریانس مترن به خوبی بازیابی شده‌اند ولی پارامتر انحراف معیار جمله نوفه کم برآورد شده است.

جدول ۴.۴: برآوردهای میانه پسینی به همراه فواصل باورمند بیزی ۹۵ درصد

پارامتر	مقدار واقعی	برآورد	فاصله باورمندپذیری	
			U	L
دامنه	۱/۵	۱/۷۸۴	۲/۶۹۹	۱/۲۶۶
واریانس	۲	۱/۹۷۴	۲/۳۲۵	۱/۷۰۲
انحراف معیار نوفه	۰/۱	۰/۰۰۹	۰/۰۲۸	۰/۰۰۴

۸.۴ نتیجه‌گیری

هدف اصلی این پایان‌نامه، ارایه و تشریح روش‌های نوین مدل‌بندی سه بعدی داده‌های زمین‌آماري به‌ویژه در زمینه داده‌های ژئومکانیکی مخازن نفتی است. برای این کار رهیافت‌های مختلفی را مطرح کردیم که می‌توان به صورت زیر خلاصه کرد:

۱. روش کریگینگ معمولی در حالت کلاسیک که در فصل دوم به آن پرداخته شد و بر روی یک مجموعه داده از یک میدان نفتی پیاده‌سازی شد.

۲. روش‌های کریگینگ میانه پالایش و عام با مدل‌بندی اسپلاین که در فصل سوم مطرح شدند و با دو مثال شبیه‌سازی عملکرد آن‌ها تشریح و با هم مقایسه شدند. روش میانه پالایش در مسایل سه بعدی برای اولین بار در این پایان‌نامه ارایه شده است.

۳. روش کریگینگ بیزی تقریبی با استفاده از رهیافت ترکیبی INLA+SPDE که در این زمینه نیز روش تازه‌ای است. عملکرد روش نیز در یک مثال شبیه‌سازی بررسی شد.

این مجموعه نشان می‌دهد که در این زمینه کارهای پژوهشی زیادی صورت نگرفته است و با توجه به کاربردهای جذاب هم از جنبه کاربردی و هم مالی، تمرکز بر روی آن‌ها می‌تواند بسیار مفید باشد.

مراجع

- [۱] حسنی پاک ع.آ. (۱۳۷۷)، زمین‌آمار (ژئواستاتیک)، چاپ اول، موسسه انتشارات و چاپ دانشگاه تهران، تهران.
- [۲] محمدزاده م. (۱۳۹۸)، آمار فضایی و کاربردهای آن، چاپ سوم، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران.
- [3] Abdel-Aal, H. K. and Alsahlawi, M. A. (2014), *Petroleum Economics and Engineering*, CRC Press, Boca Raton.
- [4] Bartlett, M. S. (1938), The approximate recovery of information from replicated field experiments with large blocks, *Journal of Agricultural Science (Cambridge)*, **28**, 418-427.
- [5] Bartlett, M. S. (1978), Nearest neighbour models in the analysis of field experiments, *Journal of the Royal Statistical Society: Series B (Methodological)*, **40**(2), 147-158.
- [6] Beauchamp, K. G. (1984), *Applications of Walsh and Related Functions: with an Introduction to Sequency Theory*, Academic press, London.
- [7] Berke, O. (2001), Modified median polish kriging and its application to the Wolfcamp-Aquifer data, *Environmetrics: The Official Journal of the International Environmetrics Society*, **12**(8), 731-748.
- [8] Besag, J. (1986), On the statistical analysis of dirty pictures, *Journal of the Royal Statistical Society: Series B (Methodological)*, **48**(3), 259-279.
- [9] Blangiardo, M. and Cameletti, M. (2015), *Spatial and Spatio-Temporal Bayesian Models with R-INLA*, John Wiley and Sons, UK, West Sussex.
- [10] Box, G. E. and Wilson, K. B. (1951), On the experimental attainment of optimum conditions, *Journal of the Royal Statistical Society: Series B (Methodological)*, **13**(1), 1-38.

- [11] Box, G. E. P., Hunter, W. G. and Hunter, J. S. (1978), *Statistics for Experimenters an Introduction to Design, Data Analysis and Model Building*, John Wiley and Sons, New York.
- [12] Brenner, S. and Scott, R. (2007), *The Mathematical Theory of Finite Element Methods*, Springer Science and Business Media, New York.
- [13] Brooks, S. Gelman, A., Jones, G. and Meng, X. L. (2011), *Handbook of Markov Chain Monte Carlo*, CRC Press, Taylor and Francis, Boca Raton.
- [14] Casella, G. and Berger, R. L. (2002), *Statistical Inference*, Brooks/Cole Cengage Learning, CA, Belmont.
- [15] Cheng, A. H. D. and Cheng, D. T. (2005), Heritage and early history of the boundary element method, *Engineering Analysis with Boundary Elements*, **29**(3), 268-302.
- [16] Chilés, J. P. and Delfiner, P. (2009), *Geostatistics: Modeling Spatial Uncertainty*, Second Edition, John Wiley and Sons, New Jersey.
- [17] Cressie, N. (1993), *Statistics for Spatial Data*, John Wiley and Sons, New York.
- [18] Cressie, N. (2015), *Statistics for Spatial Data*, John Wiley and Sons, New York.
- [19] Cressie, N. and Hawkins, D. M. (1980), Robust estimation of the variogram: I, *Journal of the International Association for Mathematical Geology*, **12**(2), 115-125.
- [20] Cressie, N. and Wikle, C. K. (2015), *Statistics for Spatio-Temporal Data*, John Wiley and Sons, New York.
- [21] Cressie, N. and Glonek, G. (1984), Median based covariogram estimators reduce bias, *Statistics and Probability Letters*, **2**(5), 299-304.
- [22] Dempster, A. P. (1974), The direct use of likelihood for significance testing , *Proceedings of Conference on Foundational Questions in Statistical Inference*, Department of Theoretical Statistics, University of Aarhus, 335-352.
- [23] Diggle, P. J. and Ribeiro, P. J.(2007), *Model based Geostatistics*, Springer series in statistics, New York.
- [24] Diggle, P. J. and Giorgi, E. (2019), *Model-based Geostatistics for Global Public Health: Methods and Applications*, CRC Press, Boca Raton.

- [25] Eriksson, M. and Siska, P. P. (2000), Understanding anisotropy computations, *Mathematical Geology*, **32**(6), 683-700.
- [26] Emerson, J. D. and Hoaglin, D. C. (2000), Analysis of two-way tables by medians, *Understanding Robust and Exploratory Data Analysis*, Edited by David C. Hoaglin, Frederick Mosteller and John W. Tukey.
- [27] Fahrmeir, L. and Tutz, G. (2001), *Multivariate Statistical Modelling based on Generalized Linear Models*, 2nd edn, Springer, Berlin.
- [28] Fisher, R. A. (1935), *The Design of Experiments*, Oliver and Boyd, Edinburgh.
- [29] Fischer, M. M. and Wang, J. (2011), *Spatial Data Analysis: Models, Methods and Techniques*, Springer Science , Berlin.
- [30] Friedman, J., Hastie, T. and Tibshirani, R. (2001), *The Elements of Statistical Learning*, Springer series in statistics, New York.
- [31] Gelman, A., Carlin, J. B., Stern, H. S., Dunson, D. B., Vehtari, A. and Rubin, D. B. (2013), *Bayesian Data Analysis*, Third Edition, CRC press, Hoboken .
- [32] Gneiting, T. and Raftery, A. E. (2007), Strictly proper scoring rules, prediction, and estimation, *Journal of the American Statistical Association*, **102**(477), 359-378.
- [33] Gómez-Rubio, V. (2020), *Bayesian Inference with INLA*, CRC Press, Boca Raton.
- [34] Gilks, W. R., Richardson, S. S. and Spiegelhalter, D. (1996), *Markov Chain Monte Carlo in Practice*, Chapman k Hall/CRC, London.
- [35] Gomez, M. and Hazen, K. (1970), *Evaluating Sulfur and Ash distribution in Coal Seams by Statistical Response Surface Regression Analysis*, US Department of Interior, Bureau of Mines.
- [36] Halley, E. (1686), An historical account of the trade winds, and monsoons, observable in the seas between and near the tropicks; with an attempt to assign the phisical cause of said winds, *Philosophical Transactions*, **183**, 153-168.
- [37] Hastings, W. K. (1970), Monte Carlo sampling methods using Markov chains and their applications, *Biometrika*, **57**(1), 97-109.

- [38] Hrafnkelsson, B., Siegert, S. Huser, R., Bakka, H. and Jóhannesson, Á. V. (2020), Max-and-Smooth: a two-step approach for approximate Bayesian inference in latent Gaussian models, *Bayesian Analysis*.
- [39] Hsiao, C. K., Huang, S. Y. and Chang, C. W. (2004), Bayesian marginal inference via candidate's formula, *Statistics and Computing*, **14**(1), 59-66.
- [40] Hubin, A. and Storvik, G. (2016), Estimating the marginal likelihood with Integrated nested Laplace approximation (INLA), *ArXiv Preprint ArXiv*, 1611.01450.
- [41] Lindgren, F., Rue, H. and Lindström, J. (2011), An explicit link between Gaussian fields and Gaussian Markov random fields: the stochastic partial differential equation approach, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, **73**(4), 423-498.
- [42] Lindgren, F. and Rue, H. (2015), Bayesian spatial modelling with R-INLA, *Journal of Statistical Software*, **63**(19), 1-25.
- [43] Ludwig, G., Zhu, J., Reyes, P., Chen, C. S. and Conley, S. P. (2020), On spline-based approaches to spatial linear regression for geostatistical data, *Environmental and Ecological Statistics*, 1-28.
- [44] Mardia, K. V., Kent, J.T. and Bibby, J. M. (1989), *Multivariate Analysis Academic Press*, London.
- [45] Martins, T. G., Simpson, D., Lindgren, F. and Rue, H. (2013), Bayesian computing with INLA: new features, *Computational Statistics and Data Analysis*, **67**, 68-83.
- [46] Martínez, W. A., Melo, C. E. and Melo, O. O. (2017). Median Polish Kriging for space-time analysis of precipitation. *Spatial Statistics*, **19**, 1-20.
- [47] Matheron, G. (1971), *The Theory of Regionalized Variables and its Applications*, École National Supérieure des Mines, Paris.
- [48] Metropolis, N., Rosenbluth, A. W., Rosenbluth, M. N., Teller, A. H. and Teller, E. (1953), Equation of state calculations by fast computing machines, *The journal of Chemical Physics*, **21**(6), 1087-1092.
- [49] Montero, J. M., Fernández-Avilés, G. and Mateu, J. (2015), Spatial and Spatio-Temporal Geostatistical Modeling and Kriging, John Wiley and Sons, West Sussex.

- [50] Montgomery, D. C. (2000), *Design and Analysis of Experiments*, 5th Ed, Wiley, New York.
- [51] Mood, A. M., Graybill, F. A. and Boes, D. C. (2017), *Introduction to the Theory of Statistics*, Chennai: McGraw Hill Education (India), Pvt. Ltd.
- [52] Mosteller, F. and Tukey, J. W. (1977), *Data Analysis and Regression: A Second Course in Statistics*, Addison-Wesley, Reading.
- [53] Nazari Ostad, M., Niri, M. E. and Darjani, M. (2018), 3D modeling of geomechanical elastic properties in a carbonate-sandstone reservoir: a comparative study of geo-statistical co-simulation methods, *Journal of Geophysics and Engineering*, **15**(4), 1419-1431.
- [54] Oliver, M. A. and Webster, R. (2015), *Basic Steps in Geostatistics: the Variogram and Kriging*, Springer International Publishing, New York.
- [55] Papadakis, J. S. (1937), Méthode statistique pour des expériences sur champ. *Bull. Inst. Amel. Plantes a Salonique*, **23**, 13-29.
- [56] Pawitan, Y. (2001), *In all Likelihood: Statistical Modelling and Inference Using Likelihood*, Oxford University Press.
- [57] Ripley, B. (1981), *Spatial Statistics*, Wiley, NewYork.
- [58] Rue, H. and Held, L. (2005), *Gaussian Markov Random Fields: Theory and Applications*, CRC press, Boca Raton.
- [59] Rue, H. and Martino, S. (2007) Approximate Bayesian inference for hierarchical Gaussian Markov random fields models, *J. Statist. Planng Inf*, **137**, 3177–3192.
- [60] Rue, H., Martino, S. and Chopin, N. (2009), Approximate Bayesian inference for latent Gaussian models by using integrated nested Laplace approximations, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, **71**(2), 319-392.
- [61] Rue, H., Riebler, A., Sørbye, S. H., Illian, J. B., Simpson, D. P. and Lindgren, F. K. (2017), Bayesian computing with INLA: a review, *Annual Review of Statistics and Its Application*, **4**, 395-421.

- [62] Sanchez, S. M. and Sanchez, P. J. (2005), Very large fractional factorial and central composite designs, *ACM Transactions on Modeling and Computer Simulation (TOMACS)*, **15**(4), 362-377.
- [63] Sarma, D. D. (2010), *Geostatistics with Applications in Earth Sciences*, Second Edition, Springer Science and Business Media, New York.
- [64] Schoenberg, I. J. (1946), Contributions to the problem of approximation of equidistant data by analytic functions. Part B. On the problem of osculatory interpolation. A second class of analytic approximation formulae, *Quarterly of Applied Mathematics*, **4**(2), 112-141.
- [65] Smith, H. (1938), An empirical law describing heterogeneity in the yields of agricultural crops, *Journal of Agricultural Science (Cambridge)*, **28**, 1-23.
- [66] Spiegelhalter, D. J., Best, N. G., Carlin, B. P. and Van Der Linde, A. (2002), Bayesian measures of model complexity and fit, *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, **64**(4), 583-639.
- [67] Stein, M. L. (1999), *Interpolation of Spatial Data: Some Theory for Kriging*, Springer, New York.
- [68] Stroup, W. W. (2012), *Generalized Linear Mixed Models: Modern Concepts, Methods and Applications*, CRC press, Hoboken.
- [69] Student, (1907), On the error of counting with a haemocytometer, *Biometrika*, **5**, 351-360.
- [70] Takahashi, K., Fagan, J. and Chin, M. S. (1973), Formation of a sparse bus impedance matrix and its application to short circuit study, *In Proc. 8th PICA Conference*, 4-6.
- [71] Tukey, J. W. (1970), *Exploratory Data Analysis*, Addison-Wesley.
- [72] Tukey, J. W. (1977), *Exploratory Data Analysis*, Addison-Wesley Company.
- [73] Velleman, P. F. and Hoaglin, D. C. (1981), *Applications, Basics, and Computing of Exploratory Data Analysis*, Scituate: Duxbury Press, New York.
- [74] Wang, X., Yue, Y. R. and Faraway, J. J. (2018), *Bayesian Regression Modeling with INLA*, CRC Press.

- [75] Whittle, P. (1954), On stationary processes in the plane, *Biometrika*, **41**(3-4) 434-449.
- [76] Whittle, P. (1963), Stochastic-processes in several dimensions, *Bulletin of the International Statistical Institute*, **40**(2), 974-994.
- [77] Wilkinson, G. N., Eckert, S. R., Hancock, T. W. and Mayo, O. (1983), Nearest neighbor (NN) analysis with field experiments, *Journal of the Royal Statistical Society B*, **45**, 151-178.
- [78] Yates, F. (1938), The comparative advantages of systematic and randomized arrangements in the design of agricultural and biological experiments, *Biometrika*, **30**, 444-466.
- [79] Zhang, R. Czado, C. and Sigloch, K. (2016), Bayesian spatial modelling for high dimensional seismic inverse problems, *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, **65**(2), 187-213.
- [80] Zimmerman, D. L. (1993), Another look at anisotropy in geostatistics, *Mathematical Geology*, **25**(4), 453-470.
- [81] Zoback, M. D. (2010), *Reservoir Geomechanics*, Cambridge University Press.

Abstract

The use of mechanical rock properties to study underground spaces is called geo-mechanic. Geo-mechanic is one of the most important branches of science in exploring and drilling underground reservoirs, especially oil reservoirs. Most of the data obtained in geo-mechanical studies are geostatistical data. Modeling and analysis of such data require spatial statistics methods. However, much of the geostatistical data in geo-mechanical studies are three-dimensional, making it challenging to develop spatial statistical models and efficient computational methods for their analysis. In this thesis, we propose two new methods for analyzing three-dimensional spatial random fields, for the first time, to address the mentioned challenges. First, we introduce a nonparametric method called three-dimensional median polish to estimate the trend function of a spatial random field. Employing this method, we also develop a median polish kriging method. We also compare the proposed median polish kriging's performance with universal kriging, in which the trend function is estimated using the spline method. The performance is assessed using both simulated and real examples of uniaxial compressive unconfined compressive strength in one of Iran's oil reservoirs. Second, using a Bayesian approach, called Integrated Nested Laplace Approximation (INLA), and its combination with Stochastic Partial Differential Equations (SPDE) method, we will develop the analysis of three-dimensional continuous indexed spatial random fields.

Keywords: Geostatistics, Bayesian inference, three-dimensional median polish Kriging, Integrated Nested Laplace Approximation, INLA+SPDE approach.



Faculty Of Mathematical Sciences

MSc Thesis in Mathematical Statistics

**Bayesian 3d geo-mechanical modeling by
using geostatistical models: a case study of
an Iranian oil field**

By: Fatane Fakhri

Supervisor:

Dr. Hossein Baghishani

Advisor:

Dr. Behzad Tokhmechi

November 2020