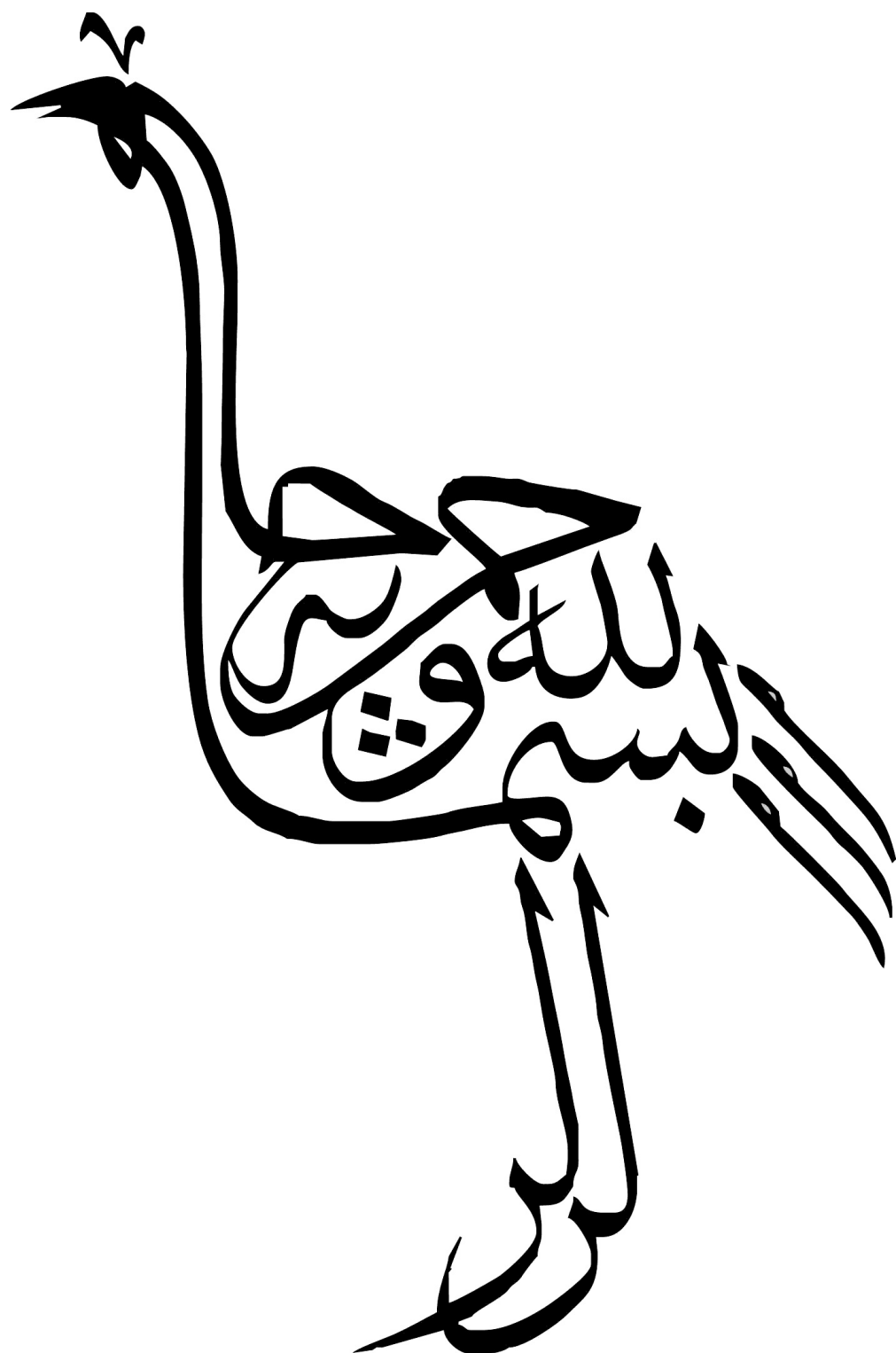






دانشکده علوم ریاضی

رساله دکتری

تحلیل بیزی داده‌های رسته‌ای با ساختارهای فضایی و فضایی-زمانی

نگارنده: لیلا برموده

استاد راهنما

دکتر حسین باغیشنی

اسفند ۱۴۰۰

شماره: ۲۵۱-۲۱۲-۲۱
تاریخ: ۱۵/۱۲/۱۴۰۰
ویرایش:

باسمه تعالی



مدیریت تحصیلات تکمیلی

پیوست شماره ۲

دانشکده: ریاضی

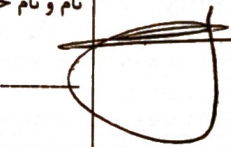
گروه: آمار

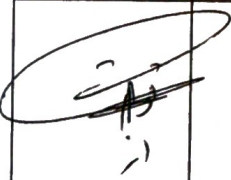
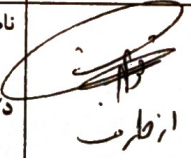
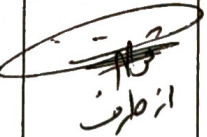
رساله دکتری خانم لیلا برموده

شماره دانشجویی: ۹۵۰۰۰۸۵

تحت عنوان: تحلیل بیزی داده های رسته ای با ساختارهای فضایی و فضایی - زمان

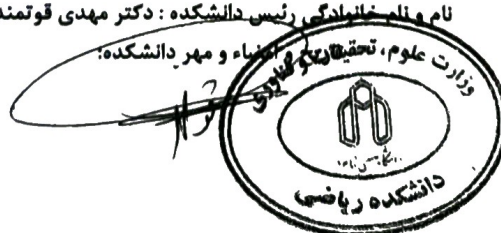
در تاریخ ۱۴۰۰/۱۲/۰۴ توسط کمیته تخصصی زیر جهت اخذ مدرک رساله دکتری ارزیابی گردید و با درجه عالی مورد پذیرش قرار گرفت.

امضاء	اساتید مشاور	امضاء	اساتید راهنما
-----	نام و نام خانوادگی:		نام و نام خانوادگی: دکتر حسین باغیشنی
-----	نام و نام خانوادگی: -----	-----	نام و نام خانوادگی: -----

امضاء	نماینده تحصیلات تکمیلی	امضاء	اساتید داور
	نام و نام خانوادگی: دکتر مهدی قوتمند		نام و نام خانوادگی: دکتر احمد نزاقتی رضازاده
			نام و نام خانوادگی: دکتر افشین فلاح
			نام و نام خانوادگی: دکتر امید کریمی

نام و نام خانوادگی رئیس دانشکده: دکتر مهدی قوتمند

امضاء و مهر دانشکده:



پیشکشی ناقابل به شهید مدافع حرم
مهندس مصطفی زال نژاد
امیدوارم که این رساله از عهده پیشکش شدن
به حضورش برآید.

حمد خالق، نعت مخلوق

حمد مخصوص خداست. حمدی که حدش بی نهایت و حسابش بی شمار و مقدارش نامتناهی و زمانش نامنقطع باشد. خدایا تو را شکر می کنم که از همان آغاز کودکی ام مرا با درد آشنا کردی تا به ارزش کیمیایی درد پی ببرم و ناخالصی های وجودم را در آتش درد بسوزم و هنگام راه رفتن بر روی زمین خاطرم آرام باشد و موجودیت وجود خود را احساس کنم. در آغاز، سپاس را مخصوص وجود پدر و مادری می دانم که صبورانه مرا در این راه یاری نموده و مرا از سرچشمه محبت خویش بهره مند ساختند. همچنین از زحمات دکتر حسین باغیشنی که راهنمایی این رساله را برعهده داشتند تشکر و قدر دانی می نمایم.

در پایان نیز وظیفه خود می بینم از برادر عزیزم که پیوسته مشوق من برای پیمودن مسیر بوده است، تشکر نمایم.

از داوران محترم که داوری این رساله را تقبل فرموده اند، نیز کمال تشکر را دارم.

لیلا برموده
اسفند ۱۴۰۰

تعهد نامه

این جانب لیلا برموده دانشجوی دکتری رشته آمار در دانشکده علوم ریاضی دانشگاه صنعتی شاهرود، نویسنده رساله با عنوان تحلیل بیزی داده‌های رسته‌ای با ساختارهای فضایی و فضایی-زمانی، تحت راهنمایی آقای دکتر حسین باغیشنی متعهد می‌شوم:

- تحقیقات در این رساله توسط این جانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش‌های دیگر پژوهش‌گران، به مرجع مورد استفاده استناد شده است.
- مطالب این رساله، تا کنون توسط خود، یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ‌جا 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 به‌دست آوردن نتایج اصلی رساله تاثیرگذار بوده‌اند، در مقالات مستخرج از رساله رعایت می‌شود.
- در تمام مراحل انجام این رساله، در مواردی که از موجود زنده (یا بافت‌های آن‌ها) استفاده شده است، ضوابط و اصول اخلاقی رعایت شده است.
- در تمام مراحل انجام این رساله، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته (یا استفاده شده است)، اصل رازداری و اصول اخلاق انسانی رعایت شده است.

لیلا برموده

اسفند ۱۴۰۰

مالکیت نتایج و حق نشر

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

چکیده

متغیر رسته‌ای، متغیری است که مشاهدات را گروه‌بندی می‌کند. برای مثال، می‌توان گروه‌بندی افراد با ویژگی‌های متفاوت و مشترک بر اساس گروه خونی را نام برد. امروزه گسترش علوم مختلف، تحلیل کارای داده‌های رسته‌ای را ضروری کرده است، به‌طوری که شاخه‌های مختلفی از علوم و کاربردهای آن‌ها به دنبال راهی برای تحلیل دقیق‌تر این نوع داده‌ها هستند. در آمار کلاسیک روش‌های متعددی برای تحلیل داده‌های رسته‌ای ارائه شده‌اند. در رهیافت بیزی تحلیل و استنباط داده‌های رسته‌ای نیازمند استفاده از روش‌های مبتنی بر نمونه‌گیری مانند الگوریتم‌های MCMC است. اما این روش‌ها برای داده‌های رسته‌ای با ساختارهای وابستگی پیچیده از کارایی لازم برخوردار نیستند. رهیافت جانشین در این رساله، استفاده از روش تقریبی بیزی لاپلاس آشیانی جمع‌بسته (INLA) است که از مشکلات عمده الگوریتم‌های MCMC در استنباط بیزی مدل‌های فضایی برحذر است. برای توسعه مدل بیزی مناسب با مشکل ناشناسایی ذاتی در مدل‌های چندجمله‌ای روبرو هستیم که برای رفع آن، از دو نوع قید برای پارامترهای مدل استفاده و کارایی آن‌ها را با هم مقایسه می‌کنیم. همچنین از تبدیل چندجمله‌ای-پواسن استفاده خواهیم کرد تا امکان برازش مدل در چارچوب مدل‌های گاوسی پنهان فراهم شود. به‌عنوان یک مدل و راه‌چاره‌ای دیگر، با توجه به نوع پاسخ، از مدل لوجیت تکی‌شده استفاده می‌کنیم. سرانجام کارایی مدل‌های پیشنهادی را به کمک مثال‌های شبیه‌سازی و واقعی در هر کدام از روش‌ها به‌طور جداگانه مورد بررسی و ارزیابی قرار می‌دهیم.

واژه‌های کلیدی: استنباط بیزی، شناسایی‌پذیری، مدل چندجمله‌ای فضایی، مدل چندجمله‌ای کسری، لوجیت تکی‌شده، تبدیل چندجمله‌ای-پواسن، ساختار وابستگی فضایی-زمانی، روش INLA.

پیش‌گفتار

توسعه و پیشرفت‌های متعدد در حوزه‌های مختلف کاربردی، گردآوری و تحلیل داده‌ها را بیش از پیش مهم و حیاتی کرده است. به خوبی می‌توان دید که صنایعی همچون گردشگری و حمل و نقل و علوم مانند علوم اجتماعی، زمین‌فیزیک و پزشکی به شدت نیازمند تحلیل داده‌ها هستند. خوشبختانه با پیشرفت ابزارهای ثبت داده‌ها و همچنین ایجاد فضاهای کافی برای ذخیره آن‌ها، تا حدی از چالش‌های پیش رو در زمینه گردآوری و ثبت داده‌ها کاسته شده است. امروزه حتی می‌توان از به‌روز رسانی دقیقه‌ای و ثانیه‌ای داده‌های با حجم بالا استفاده کرد. گردآوری داده‌ها و سازماندهی آن‌ها تا زمانی که به تصمیم‌گیری و نتیجه خاصی منجر نشود، کار چندان مفیدی نخواهد بود. در واقع، داده‌ها زمینه‌ای را فراهم می‌آورند تا با پردازش و تحلیل آن‌ها درباره پدیده تحت بررسی و قضاوت، تصمیم‌گیری و برنامه‌ریزی کنیم. در این میان، نمی‌توان از اهمیت داده‌های رسته‌ای در تحلیل داده‌ها غافل شد زیرا علاوه بر ذات طبیعی رده‌ای بودن بسیاری از داده‌ها، همواره با داده‌های روبرو هستیم که در تعیین و گردآوری آن‌ها به آزمون و خطا نیاز داریم که خود موجب رده‌ای شدن داده‌ها می‌شود. برای مثال، فرض کنید بخواهیم میزان دوز مؤثر یک دارو را بیابیم، برای این کار باید در ابتدا رده‌هایی از آن را مورد بررسی قرار دهیم.

از طرف دیگر، پیشرفت در علم داده‌ها موجب شده است تا هم‌زمان داده‌های مربوط به مکان و نواحی مختلف در زمان‌های متوالی و مختلف در دسترس باشند. برای بررسی بهتر و جامع‌تر تفاوت‌های منطقه‌ای داده‌ها و کشف روند تغییرات و حرکت آن‌ها باید اثرات مربوط به نواحی مورد مطالعه و زمان آن‌ها را نیز در نظر بگیریم، زیرا تجربه محققان در طول سالیان دراز نشان داده است که داده‌ها در نواحی مختلف و در بازه‌های زمانی مختلف از تفاوت‌های چشمگیری برخوردارند. اگرچه در نظر گرفتن آن‌ها به میزان زیادی بر پیچیدگی مدل‌ها خواهد افزود.

برای تحلیل و بررسی دقیق داده‌ها نیازمند یافتن و به‌کارگیری مدل‌هایی هستیم که بتواند همه اطلاعات موجود در داده‌ها را به روشنی بیان کند؛ زیرا یک مدل خوب باید علاوه بر توانایی توصیف داده‌ها، قدرت پیشگویی بالایی داشته باشد. افزودن هم‌زمان پیچیدگی مدل با در نظر گرفتن اثرات مربوط به مکان داده‌ها و زمان گردآوری آن‌ها، نباید ما را از مسئله قابل شناسایی بودن مدل غافل کند. به عبارت دیگر، مدل‌ها باید توانایی تشخیص پارامترهای مدل را به خوبی داشته باشند.

با توجه به نکاتی که ذکر شد، در این رساله هدف اصلی مدل‌بندی بیزی داده‌های رسته‌ای با تمرکز بر داده‌های (اسمی) چندجمله‌ای، به دلیل اهمیت و فراوانی آن‌ها در علوم مختلف، با حضور اثرات مربوط به مکان و زمان است. یکی از مسایل مهم مورد نظر ما در این راستا، بررسی و اعمال شرایط شناسایی‌پذیری مدل در کنار کارایی محاسباتی و آماری برازش مدل است. در استنباط بیزی معمولاً از الگوریتم‌های نمونه‌گیری MCMC برای تحلیل مدل‌های

آماري پيچيده استفاده مي کنند، اما اين الگوريتمها به ويژه براي داده هاي رسته اي چند جمله اي از چالش هاي جدی مانند همگرایی ضعيف زنجير مارکوف رنج مي برند. از طرف ديگر، افزودن اثرات مربوط به مکان و زمان داده ها به وضوح بر اين دشواري مي افزايد. براي دوري از مشکلات عمده الگوريتم هاي MCMC در استنباط بيزي مدل هاي فضايي و فضايي-زمانی، رهيافت پيشنهادي در اين رساله استفاده از روش تقريبي بيزي لاپلاس آشياني جمع بسته (INLA) است. اهداف اين رساله را مي توان به صورت زير خلاصه کرد:

- افزايش قدرت شناسايي پذيري مدل به طوري که دقت برآوردها افزايش يابد.
 - افزايش کارايي و سرعت در برازش مدل.
 - توسعه مدل هاي رسته اي با در نظر گرفتن اثرات ذاتي فضايي و فضايي-زمانی در چارچوب بيزي.
 - توسعه مدل رگرسيون جمعي-ساختاري براي داده هاي رسته اي در چارچوب بيزي.
 - مقايسه مدل پيشنهادي با ساير مدل هاي متداول براي تحليل داده هاي رسته اي.
- با توجه به نکاتي که در بالا ذکر شد، محتوای فصل هاي رساله را مي توان به صورت زير خلاصه کرد:
- فصل اول** به معرفي خلاصه اي از داده هاي رسته اي و مدل هاي متداول به کار رفته در تحليل آن ها مي پردازد. ساختارهاي وابستگي خاص مانند ساختار فضايي و فضايي-زمانی در اين فصل معرفي مي شوند.
- فصل دوم** رهيافت INLA را به طور خلاصه شرح مي دهد.
- فصل سوم** به تشریح کامل مفهوم شناسايي پذيري در حالت کلی و به طور خاص در رويکرد بيزي مي پردازد. همچنين مشکلات شناسايي پذيري را معرفي و راه حل هاي آن را به همراه نتايج نظري ارائه مي کند.
- فصل چهارم** مدل چند جمله اي با ساختار جمعي-رگرسيوني با در نظر گرفتن قيد شناسايي در حالت فضايي را ارزيابي کرده و به کمک مثال هاي شبیه سازی و داده واقعي نتايج حاصل را با مدل متداول چند جمله اي کسري مقايسه و تحليل مي کند.
- فصل پنجم** مدل لجيت تکی شده براي رفع مشکل شناسايي پذيري مدل چند جمله اي را در حالي بررسی مي کند که داده ها فراواني ندارند و اثر فضايي از نوع زمين آماری است.
- فصل ششم** مسئله شناسايي پذيري را براي مدل چند جمله اي جمعي-ساختاري فضايي-زمانی بررسی مي کند.

فهرست مقالات مستخرج از رساله

1. Barmoudeh, L., Baghishani, H., Martino, S. (2022). Bayesian spatial analysis of crash severity data with the INLA approach: assessment of different identification constraints. Accident Analysis and Prevention. To appear.
2. Barmoudeh, L., Baghishani, H. (2022). Spatial Bayesian individualized logit models with application to water quality data analysis. In preparation.

فهرست مطالب

ق فهرست تصاویر

ش فهرست جداول

۱	تحلیل داده‌های رسته‌ای	۱
۱	۱.۱ مقدمه	۱
۲	۲.۱ انواع داده‌های رسته‌ای	۲
۲	۱.۲.۱ متغیرهای رسته‌ای اسمی	۲
۳	۲.۲.۱ متغیرهای رسته‌ای ترتیبی	۳
۳	۳.۱ توزیع‌های داده‌های رسته‌ای	۳
۳	۱.۳.۱ توزیع دوجمله‌ای	۳
۵	۲.۳.۱ توزیع چندجمله‌ای	۵
۵	۳.۳.۱ توزیع پواسن	۵
۶	۴.۱ مدل‌های رگرسیونی برای داده‌های رسته‌ای	۶
۶	۱.۴.۱ خانواده توزیع‌های نمایی	۶
۷	۲.۴.۱ مدل‌های خطی تعمیم‌یافته	۷
۱۵	۵.۱ ساختارهای وابستگی	۱۵
۱۵	۱.۵.۱ مدل‌های آمیخته خطی تعمیم‌یافته	۱۵
۱۵	۲.۵.۱ مدل آمیخته جمعی-ساختاری تعمیم‌یافته	۱۵
۱۷	۶.۱ ساختارهای وابستگی خاص	۱۷
۱۸	۱.۶.۱ ساختار وابستگی زمانی	۱۸
۱۹	۲.۶.۱ ساختار وابستگی فضایی	۱۹
۲۱	۳.۶.۱ وابستگی فضایی-زمانی	۲۱
۲۳	۲ تقریب لاپلاس آشیانه‌ای جمع‌بسته	۲۳
۲۳	۱.۲ مقدمه‌ای بر INLA	۲۳

۲۵	تقریب لاپلاس	۲.۲
۲۶	مدل‌های گوسی پنهان	۱.۲.۲
۲۷	میدان‌های تصادفی گوسی مارکوفی	۲.۲.۲
۲۸	مراحل روش تقریب INLA	۳.۲
۳۲	مثال‌های بیشتر برای مقایسه	۴.۲
۳۵	۳ شناسایی‌پذیری	
۳۵	مقدمه	۱.۳
۳۶	شناسایی‌پذیری مدل	۲.۳
۳۹	شناسایی‌پذیری و اهمیت آن در استنباط بیزی	۳.۳
۴۰	حل مسئله شناسایی‌پذیری	۴.۳
۴۰	تاریخچه استفاده از قیدها در آمار	۱.۴.۳
۴۱	مدل نرمال با ویژگی ناشناسایی مشابه با مدل لوجیت چندجمله‌ای	۵.۳
۴۷	۴ مدل چندجمله‌ای جمعی-ساختاری با در نظر گرفتن قید شناسایی	
۴۸	تعریف مدل	۱.۴
۴۸	رگرسیون جمعی-ساختاری برای داده‌های رسته‌ای اسمی	۱.۱.۴
۴۹	مدل‌بندی ساختار وابستگی فضایی	۲.۱.۴
۵۴	شناسایی‌پذیری در مدل‌های چندجمله‌ای	۳.۱.۴
۵۶	تبدیل چندجمله‌ای-پواسن	۴.۱.۴
۵۸	تبدیل چندجمله‌ای-پواسن بیزی	۵.۱.۴
۵۹	مدل چندجمله‌ای کسری	۶.۱.۴
۶۰	مطالعه شبیه‌سازی	۲.۴
۶۳	مثال کاربردی: شدت تصادفات جاده‌ای	۳.۴
۶۹	بررسی مدل‌ها از نظر قدرت پیشگویی	۱.۳.۴
۷۱	انتخاب مدل	۲.۳.۴
۷۳	خلاصه فصل و نتیجه‌گیری	۴.۴
۷۹	۵ مدل لوجیت تکی‌شده برای داده‌های چندجمله‌ای فضایی زمین‌آماري	
۷۹	مدل لوجیت تکی‌شده	۱.۵
۸۰	توزیع مجانبی برآورد درست‌نمایی ماکسیمم	۱.۱.۵
۸۲	همگرایی برآوردگر بیزی مدل تکی‌شده	۲.۱.۵
۸۳	مدل‌بندی وابستگی فضایی زمین‌آماري	۲.۵
۸۷	پارامتربندی و پیشین‌ها	۱.۲.۵

۳.۵	مقایسه روش‌های لوجیت تکی‌شده و چندجمله‌ای به کمک شبیه‌سازی و
۸۷	مثال واقعی
۸۷	۱.۳.۵ مطالعه شبیه‌سازی
۸۹	۴.۵ مثال کاربردی: کیفیت آب‌های زیرزمینی
۹۲	۱.۴.۵ مدل‌بندی و برآورد پارامتر
۹۲	۵.۵ خلاصه و نتیجه‌گیری
۹۷	۶ مدل‌بندی داده‌های چندجمله‌ای فضایی-زمانی
۹۷	۱.۶ اثر فضایی-زمانی
۹۸	۱.۱.۶ اثرات متقابل ذاتی فضایی-زمانی
۹۹	۲.۱.۶ اثرات اصلی
۹۹	۳.۱.۶ اثرات متقابل فضایی-زمانی
۱۰۰	۴.۱.۶ قیدها و فضای صفر
۱۰۲	۲.۶ راه‌های دوری از قید مجموع مساوی با صفر در اثرات متقابل
۱۰۲	۱.۲.۶ پارامتربندی متفاوت
۱۰۴	۲.۲.۶ پیاده کردن قیدها بر روی نمونه‌های توزیع پسین
۱۰۵	۳.۶ مطالعه شبیه‌سازی
۱۰۷	۴.۶ برازش مدل با اثر فضایی-زمانی با قید رده گوشه
۱۰۸	۱.۴.۶ مطالعه شبیه‌سازی
۱۱۰	۵.۶ خلاصه فصل و نتیجه‌گیری
۱۱۰	۶.۶ آینده پژوهش
۱۱۵	مراجع
۱۳۳	واژه‌نامه فارسی به انگلیسی
۱۳۷	واژه‌نامه انگلیسی به فارسی

فهرست تصاویر

۱۰	۱.۱	منحنی رگرسیون احتمال پیروزی به ازای یک متغیر تبیینی، در حالتی که ضریب رگرسیونی مثبت (سمت چپ) یا منفی (سمت راست) است.
۴۴	۱.۳	نسبت میانگین پسین با در نظر گرفتن قیدهای شناسایی پذیری در مدل نرمال عاملی.
	۲.۳	مقایسه دقت با در نظر گرفتن قیدهای شناسایی پذیری در مدل نرمال عاملی برای $n = 30$ و $\tau_y = 1$
۴۴		دقت‌های برآوردشده در هر رده وقتی τ_B از ۱ تا ۱۰۰٪ تغییر می‌کند، تقریباً به یکدیگر نزدیک هستند.
۴۴	۳.۳	مقایسه واریانس با در نظر گرفتن قیدهای شناسایی پذیری در مدل نرمال عاملی وقتی تعداد مشاهدات رو به افزایش است.
۵۰	۱.۴	مدل بی‌سیگ. شکل سمت چپ مربوط به داده‌های یک مثال از نواحی کشور آلمان است. شکل سمت راست ماتریس دقت تنک Q متناظر با این نواحی را نشان می‌دهد.
۵۱	۲.۴	نمونه‌ای از یک گراف متصل‌شده.
۶۲	۳.۴	میانگین پسین اثر فضایی در هر رده: نواحی سایه‌زده‌شده کران اعتبار بیزی ۹۵٪ را نشان می‌دهد. از بالا به پایین به ترتیب متعلق به رده‌های اول تا سوم است.
۶۳	۴.۴	برآورد احتمال هر رده: از بالا به پایین به ترتیب متعلق به رده‌های اول تا سوم است.
۶۴	۵.۴	چگالی‌های کناری پسین پارامتر دقت τ_u : قید مجموع مساوی با صفر (چپ)، و قید رده گوشه (راست). خط عمودی مقادیر واقعی را نشان می‌دهد.
۶۶	۶.۴	موقعیت جغرافیایی استان مازندران در نقشه ایران (قاب سمت چپ)، و نواحی مورد مطالعه (قاب سمت راست)
۶۷	۷.۴	توزیع فراوانی مربوط به رده‌های جراحات برای همه ۲۲ ناحیه استان مازندران.
۷۵	۸.۴	برآوردهای میانگین پسینی اثر سن با کران اعتبار بیزی ۹۵٪ در داده‌های تصادفات مازندران
۷۶	۹.۴	برآوردهای میانگین پسینی اثر فضایی در هر رده و برای هر مدل: از بالا به پایین مربوط به رده‌های با شدت جراحات کم تا بالا.
۷۶	۱۰.۴	برآورد انحراف استاندارد پسینی اثر فضایی در هر رده برای هر مدل: از بالا به پایین مربوط به رده‌های با شدت جراحات کم تا بالا.
۷۷	۱۱.۴	توزیع احتمال برآوردشده در هر رده برای هر مدل: از بالا به پایین مربوط به رده‌های با شدت جراحات کم تا بالا.

۱.۵	برآورد احتمال رده‌های دوم و سوم در سناریوهای نامتعادل (چپ) و متعادل (راست). ردیف بالا
۸۹	مربوط به رده دوم و ردیف پایین مربوط به رده سوم است.
۲.۵	استان گلستان: ناحیه هاشورخورده در نقشه ایران.
۳.۵	میانگین پسینی اثر فضایی برای رده اول (ردیف اول) و رده دوم (ردیف دوم) در دو مدل لوجیت‌های
۹۳	تکی‌شده (سمت چپ) و چندجمله‌ای (سمت راست).
۴.۵	انحراف استاندارد پسینی اثر فضایی برای رده اول (ردیف اول) و رده دوم (ردیف دوم) در دو مدل
۹۴	لوجیت‌های تکی‌شده (سمت چپ) و چندجمله‌ای (سمت راست).
۵.۵	برآوردهای احتمال برای رده اول (ردیف اول) و رده دوم (ردیف دوم) در دو مدل لوجیت‌های تکی‌شده
۹۵	(سمت چپ) و چندجمله‌ای (سمت راست).
۱.۶	میانگین پسین اثر زمانی (دایره‌های سبزرنگ) به همراه اثر زمانی شبیه‌سازی‌شده واقعی (خط مشکی)
۱۰۷	برای هر سه رده.
۲.۶	میانگین پسین اثر فضایی (دایره‌های قرمز رنگ) به همراه اثر فضایی شبیه‌سازی‌شده واقعی (خط
۱۰۸	مشکی) برای هر سه رده.
۳.۶	میانگین پسین اثر متقابل فضایی-زمانی (دایره‌های سبزرنگ) به همراه اثر متقابل شبیه‌سازی‌شده
۱۰۸	واقعی (خط مشکی) برای هر سه رده.
۴.۶	برآوردهای میانگین پسین اثر عرض از مبدأ با فاصله اطمینان ۹۵٪ (چپ)، برآوردهای میانگین پسین
۱۰۹	اثر ضریب رگرسیونی با فاصله اطمینان ۹۵٪ (راست). برآوردها به رنگ قرمز نمایش داده شده‌اند و
۱۰۹	مقادیر واقعی به رنگ مشکی.
۵.۶	برآوردهای میانگین پسین اثر فضایی در هر رده با فاصله اطمینان ۹۵٪. رده‌ها در شکل از سمت چپ
۱۱۰	به ترتیب اول تا سوم هستند. مقادیر واقعی به رنگ مشکی و برآوردها به رنگ قرمز هستند.
۶.۶	برآوردهای میانگین پسین اثر زمانی در هر رده با فاصله اطمینان ۹۵٪. رده‌ها در شکل از سمت چپ
۱۱۱	به ترتیب اول تا سوم هستند. مقادیر واقعی به رنگ مشکی و برآوردها به رنگ قرمز هستند.
۷.۶	برآوردهای میانگین پسین اثر فضایی-زمانی در هر رده با فاصله اطمینان ۹۵٪. رده‌ها در شکل از
۱۱۲	سمت چپ به ترتیب اول تا سوم هستند مقادیر واقعی به رنگ مشکی و برآوردها به رنگ قرمز هستند..
۸.۶	برآوردهای چگالی‌های کناری دقت برای هر اثر.
۹.۶	برآورد احتمال پسین مربوط به هر رده برای ۱۰۰ نقطه فضایی-زمانی. رده‌ها در شکل از سمت چپ
۱۱۳	به ترتیب اول تا سوم هستند.

فهرست جداول

۷	توزیع‌های پی‌آسن، دوجمله‌ای و نرمال به‌عنوان مثالی از خانواده توزیع‌های نمایی در حالت تک متغیره	۱.۱
۱۲	پارامترهای برآوردشده مدل لوجیت چندجمله‌ای در مثال سوسمارها	۲.۱
۳۲	نتایج برازش مدل رگرسیون ساده	۱.۲
۳۲	مقایسه زمان اجرا برای مدل رگرسیون ساده	۲.۲
۳۳	نمایی از ساختار داده‌های مربوط به مثال اندازه دور سر جنین	۳.۲
۳۳	برآوردهای عرض از مبدأ و شیب در مدل	۴.۲
۳۳	برآوردهای پارامترهای دقت اثرات تصادفی	۵.۲
۳۳	برآورد پارامتر دقت مشاهدات	۶.۲
۳۴	برآورد همبستگی بین شیب و عرض از مبدأ	۷.۲
	برآورد میانگین (انحراف استاندارد) پسینی برای پارامترهای مدل با استفاده از دو قید شناسایی. از روابط (۸.۴) استفاده شده است تا برآورد پارامترهای رده گوشه بازیابی شوند.	۱.۴
۶۱	میانگین (انحراف استاندارد) متغیرهای تبیینی در هر رده از پاسخ	۲.۴
۶۷	نتایج آزمون HM برای بررسی ویژگی IIA در مجموعه داده تصادف‌های مازندران. علامت * سطح	۳.۴
۶۸	معنی‌داری برابر ۱ را نشان می‌دهد.	۴.۴
	خلاصه‌ای از آماره‌های پسینی با استفاده از مدل‌های چندجمله‌ای STAR و چندجمله‌ای کسری برای داده‌های تصادف مازندران.	۷.۰
	میانگین (انحراف استاندارد) پسینی ضرایب رگرسیونی در مدل‌های لوجیت‌های تکی‌شده و چندجمله‌ای برای سناریوی نامتعادل.	۱.۵
۸۹	میانگین (انحراف استاندارد) پسینی ضرایب رگرسیونی در مدل‌های لجیت‌های تکی‌شده و چندجمله‌ای برای سناریوی متعادل.	۲.۵
۹۰	متغیرهای تبیینی موجود در مجموعه داده آب استان گلستان	۳.۵
۹۱	خلاصه‌ای از برآوردهای ضرایب رگرسیونی و پارامترهای کوواریانس مترن دو مدل برای مجموعه داده کیفیت آب استان گلستان.	۴.۵
۹۳	مقایسه میانگین و انحراف معیار برآوردشده با مقادیر شبیه‌سازی‌شده در مثال ۱.۲.۶.	۱.۶

۲.۶	برآورد میانگین (انحراف استاندارد) پسینی اثرات رگرسیونی ثابت.	۱۰۷
-----	--	-----

فصل ۱

تحلیل داده‌های رسته‌ای

۱.۱ مقدمه

تحلیل داده‌های رسته‌ای^۱ به دلیل کاربردهای بسیاری که در علوم مختلف از جمله علوم زیستی، پزشکی و علوم اجتماعی دارد، سال‌ها است که علاقه‌مندان زیادی را در این علوم و همچنین علم آمار به خود جذب کرده است. در این راستا در دهه‌های اخیر، با گسترش روش‌های استنباط آماری، کارهای متعددی در این زمینه از دانشمندان علوم مختلف به خصوص در پزشکی به چشم می‌خورد. متغیرهای رسته‌ای را می‌توان به متغیرهایی نسبت داد که فقط با تعدادی از مقادیر یا رسته اندازه‌گیری می‌شوند که با تعریف متغیر تصادفی پیوسته می‌تواند تعداد نامتناهی از مقادیر را اختیار کند، تفاوت دارد، اما گاهی حتی با داده‌های پیوسته به صورت داده‌های رسته‌ای برخورد می‌شود (دنیل، ۱۹۹۹). برای مثال، سن اگرچه به طور ذاتی یک متغیر پیوسته است، اما به دلایل کاربردی با آن مانند یک متغیر رسته‌ای برخورد می‌کنند. مثلاً ممکن است گروه‌های سنی به صورت فاصله‌های پنج ساله نمایش داده شوند. به عبارتی می‌توان گفت متغیر رسته‌ای، متغیری است که مشاهدات را گروه‌بندی می‌کند. در واقع در بسیاری از علوم با داده‌های خامی روبرو هستیم که برای تحقیق و بررسی بهتر با آن‌ها رسته‌ای برخورد می‌شود. نقطه مشترک در بین داده‌های مربوط به تولد، مرگ، ازدواج، طلاق، آموزش و پرورش، استخدام، اشتغال و مهاجرت، رسته‌ای بودن آن‌ها است که به طور گسترده در علوم

^۱Categorical

اجتماعی مورد بررسی قرار می‌گیرند. در حقیقت، بیشتر مشاهدات در تحقیقات علوم اجتماعی به‌طور رسته‌ای بررسی می‌شوند. توجه کنید حتی سوال ما با این عنوان که آیا در تحقیقات از متغیر رسته‌ای استفاده می‌شود یا نه، خود یک اندازه‌گیری رسته‌ای محسوب می‌شود.

۲.۱ انواع داده‌های رسته‌ای

بسیاری از متغیرهای رسته‌ای فقط شامل دو رسته هستند؛ مانند متغیرهایی که در آن فقط بر رسته‌های پیروزی و شکست، تاکید دارد که متغیرهای دودویی^۲ نامیده می‌شوند. در بسیاری از بخش‌های علوم مختلف، خروجی‌های حاصل از آزمایش تنها دو مقدار را می‌پذیرد. به‌عنوان مثال، می‌توان به برد و باخت یک بازی بسکتبال، اصابت یا عدم اصابت ضربه توپ توسط یک بازیکن، قبولی یا رد شدن یک دانش‌آموز در امتحان و جواب دادن و ندادن بدن یک بیمار به نوعی از داروی جدید در طول دوره درمان اشاره کرد. در همه این مثال‌ها برای راحتی کار، برآمد آزمایش را به‌صورت پیروزی و شکست بیان می‌کنند. در حالت کلی برای پیروزی پاسخ $Y = 1$ و برای شکست پاسخ $Y = 0$ را در نظر می‌گیرند. وقتی که متغیرهای رسته‌ای بیشتر از دو رسته داشته باشند، آن‌ها را به دو دسته اسمی^۳ و ترتیبی^۴ تقسیم می‌کنند.

۱.۲.۱ متغیرهای رسته‌ای اسمی

مشاهدات این نوع متغیرها مقادیر کمی ندارند و برای نامیدن یک متغیر به کار می‌روند. برای مثال

- جنسیت (زن و مرد)
- رنگ مو (قهوه‌ای، سیاه، روشن)
- محل زندگی (شمال استوا، جنوب استوا، هیچکدام)
- نظرات دانش‌آموزان درباره یک برنامه کلاسی (بی تفاوت، معمولی، عالی) و
- انتخاب یکی از چهار برند در تولید یک محصول خاص مثل ماست

از جمله این نوع داده‌ها هستند.

^۲ Binary

^۳ Nominal

^۴ Ordinal

۲.۲.۱ متغیرهای رسته‌ای ترتیبی

گونه دیگری از داده‌های رسته‌ای دارای رسته‌های ترتیبی هستند که بر اساس مقیاس ترتیبی اندازه‌گیری می‌شوند. از جمله مثال‌های این نوع داده‌ها می‌توان موارد زیر را فهرست کرد:

- بیان شرایط یک بیمار بعد از مصرف نوع جدیدی از دارو به صورت خوب، متعادل، بد و بحرانی که به ترتیب شدت را نشان می‌دهند.
- بیان میزان رضایت از خدمات یک هتل (خیلی ناراضی، ناراضی، تا حدی راضی، راضی، خیلی راضی)
- تعیین طبقه اجتماعی افراد (پایین، متوسط، بالا)
- سطح سواد افراد (راهنمایی، دبیرستان، دانشگاه) و
- رتبه‌بندی نظر یک بیننده نسبت به یک برنامه تلویزیونی (ضعیف، متوسط، خوب).

لازم به ذکر است در این رساله به علت کاربردهای متعدد و بسیاری که در طبیعت و مطالعات گذشته از داده‌های اسمی به چشم می‌خورد، تمرکز بر تحلیل داده‌های رسته‌ای اسمی است. به همین منظور، تنها در این فصل، به اختصار به تحلیل داده‌های ترتیبی اشاره خواهیم کرد. توجه داشته باشید که در این رساله برای نشان دادن بردار و ماتریس از صورت پررنگ حروف لاتین استفاده می‌شود.

۳.۱ توزیع‌های داده‌های رسته‌ای

تحلیل و استنباط داده‌ها نیازمند داشتن فرض‌هایی در مورد سازوکار تولید تصادفی داده‌ها است. برای مثال، در تحلیل داده‌های پیوسته، توزیع نرمال، نقش کلیدی را بازی می‌کند. در این بخش، سه توزیع کلیدی برای تحلیل پاسخ‌های رسته‌ای، یعنی توزیع دوجمله‌ای، توزیع چندجمله‌ای و توزیع پواسن را مرور خواهیم کرد.

۱.۳.۱ توزیع دوجمله‌ای

عموما هدف محققان از تحلیل داده‌های دودویی، برآورد و پیشگویی احتمال پیروزی یا شکست در حضور برخی متغیرهای تبیینی است. در متغیرهای دودویی واحدهای آزمایش به صورت تکی بررسی می‌شوند، به طوری که برای هر واحد تنها دو رسته پیروزی و شکست امکان‌پذیر است. چنین آزمایشی را آزمایش برنولی^۵ می‌نامند (اگرستی، ۲۰۱۳) که تنها دارای یک پارامتر

^۵Bernoulli

برای پیروزی است. تابع جرم احتمال برنولی

$$P(Y = y|\pi) = \pi^y(1 - \pi)^{1-y} \quad y = 0, 1,$$

است، که در آن $P(Y = 1) = \pi$ احتمال پیروزی را نشان می‌دهد. فرض کنید $(y_1, y_2, \dots, y_n)$ بردار مشاهدات حاصل از n آزمایش مستقل برنولی باشد. مجموع تعداد موفقیت‌ها یعنی $Y = \sum_{i=1}^n Y_i$ دارای توزیع دوجمله‌ای^۶ با تابع احتمال

$$P(Y = y) = \binom{n}{y} \pi^y (1 - \pi)^{n-y} \quad y = 0, 1, 2, \dots, n,$$

است. چون $E(Y_i) = E(Y_i^2) = \pi$ ، بنابراین امید ریاضی و واریانس توزیع دوجمله‌ای برای Y به ترتیب برابر با

$$E(Y) = n\pi \quad \text{Var}(Y) = n\pi(1 - \pi)$$

است. در حالت کلی و به‌طور خلاصه آزمایش دوجمله‌ای باید دارای ویژگی‌های زیر باشد:

۱. آزمایش شامل n آزمایش برنولی یکسان است، به این معنی که احتمال موفقیت، π ، برای هر آزمایش برنولی ثابت است.

۲. هر آزمایش دارای دو برآمد است که یکی از برآمدها را موفقیت و دیگری را شکست می‌نامند.

۳. آزمایش‌ها مستقل هستند، یعنی نتیجه یک آزمایش تأثیری در نتایج آزمایش‌های دیگر ندارد.

۴. متغیر تصادفی Y برابر با تعداد موفقیت‌های دیده‌شده در n آزمایش است.

مثال ۱.۳.۱. بسیاری از وقایع در زندگی واقعی را می‌توان با توزیع احتمال دوجمله‌ای توضیح داد. در زیر به دو کاربرد اشاره می‌کنیم.

- توزیع دوجمله‌ای به‌طور عمده در زمینه داروسازی استفاده می‌شود. برای مثال، زمانی که یک دارو برای معالجه یک بیماری خاص کشف یا اختراع می‌شود، نتایج حاصل از آزمایش با دو برآمد نمایش داده می‌شود؛ یعنی دارو در درمان بیماری مد نظر مفید است یا خیر. بنابراین با دو نتیجه موفقیت و شکست روبرو هستیم که برای مدل‌بندی موثر بودن دارو بر اساس تعداد بهبودیافته‌ها از توزیع دوجمله‌ای استفاده می‌کنیم.
- یک کاربرد توزیع دوجمله‌ای می‌تواند در مدل‌بندی تعداد پیامک‌های ناشناس دریافتی افراد باشد، زیرا در دریافت هر پیامک دو نتیجه ممکن وجود دارد: یک پیامک ناشناس است یا نیست.

^۶ Binomial

۲.۳.۱ توزیع چندجمله‌ای

حال فرض کنید برای n آزمایش مستقل، برآمدها از $J (\geq 2)$ رشته باشند. همچنین فرض کنید اگر برآمد حاصل از آزمایش i متعلق به رشته j باشد، آن گاه $Y_{ij} = 1$ و در غیر این صورت $Y_{ij} = 0$. بنابراین $y_i = (y_{i1}, y_{i2}, \dots, y_{iJ})'$ نشان‌دهنده مشاهده یک آزمایش چندجمله‌ای با $\sum_j y_{ij} = 1$ است. فرض کنید $N_j = \sum_i Y_{ij}$ تعداد برآمدهای مربوط به رشته j ام باشد. بنابراین $(N_1, N_2, \dots, N_J)'$ دارای توزیع چندجمله‌ای است. حال فرض کنید در هر بار آزمایش، احتمال این که نتیجه متعلق به رشته j ام باشد برابر با $\pi_j = P(Y_{ij} = 1)$ است. در این صورت، تابع جرم احتمال چندجمله‌ای به صورت

$$P(N_1 = n_1, \dots, N_J = n_J) = \left(\frac{n!}{n_1! \dots n_J!} \right) \pi_1^{n_1} \pi_2^{n_2} \dots \pi_J^{n_J}$$

است که در آن $\sum_j n_j = n$. با توجه به این که $n_J = n - (n_1 + n_2 + \dots + n_{J-1})$ پس توزیع $J-1$ بعدی است. توزیع دوجمله‌ای حالت خاصی از توزیع چندجمله‌ای با $J = 2$ است. میانگین، واریانس و کوواریانس برای این توزیع به صورت

$$E(N_j) = N\pi_j, \quad \text{Var}(N_j) = N\pi_j(1 - \pi_j), \quad \text{Cov}(N_j, N_k) = -N\pi_j\pi_k,$$

هستند. می‌توان نشان داد که به ازای $j = 1, \dots, J$ توزیع کناری N_j ، دوجمله‌ای است (اگرستی، ۲۰۱۳). مثال‌های زیر که از اگرستی (۲۰۰۲) برگرفته شده‌اند، برخی کاربردهای توزیع چندجمله‌ای را نشان می‌دهند.

- یک زیست‌شناس در نظر دارد تا بداند گرایش سوسمارهای موجود در یک برکه به چه نوعی از غذا بیشتر است. برای این کار، او در ابتدا انواع غذاها را رده‌بندی می‌کند. بنابراین متغیر پاسخ در این مثال، تمایل سوسمارها به انواع غذاها، دارای توزیع چندجمله‌ای است و عواملی مثل اندازه سوسمار و عوامل محیطی دیگر را می‌توان به عنوان متغیرهای تبیینی در نظر گرفت.
- در علوم مربوط به بازاریابی^۷ به طور عمده از توزیع چندجمله‌ای استفاده می‌کنند (چن و همکاران، ۲۰۰۱). با این هدف، تعدادی از برندهای خاص (مثل لوازم خانگی) را به عنوان متغیر پاسخ انتخاب کرده و عوامل تأثیرگذار مختلف را بر روی آن می‌سنجند.

۳.۳.۱ توزیع پواسن

گاهی، داده‌های شمارشی^۸ نتیجه تعداد ثابتی از آزمایش‌ها نیست. برای مثال، اگر متغیر Y تعداد افراد کشته شده در تصادفات جاده هراز شهرستان آمل در هفته آینده باشد، آن گاه حد

^۷Marketing

^۸Count data

بالایی برای تعداد n وجود ندارد. یک توزیع پرکاربرد برای این نوع متغیر، توزیع پواسن است. تابع جرم احتمال پواسن به صورت

$$P(Y = y) = \frac{e^{-\lambda} \lambda^y}{y!}, \quad y = 0, 1, \dots,$$

است. توزیع پواسن دارای میانگین و واریانس یکسان و برابر با λ است. از توزیع پواسن می‌توان برای تقریب توزیع دوجمله‌ای زمانی که اندازه n بزرگ و π کوچک باشد، نیز استفاده کرد.

۴.۱ مدل‌های رگرسیونی برای داده‌های رسته‌ای

به روشنی مدل‌بندی داده‌ها به ما کمک می‌کند تا ارزیابی کاراتری از تفسیر روند داده‌ها و تأثیرات مربوط به آن‌ها داشته باشیم. امروزه با گسترش روش‌های آماری هم‌زمان با حضور نرم‌افزارهای قوی ما را به سمت موقعیت‌هایی می‌کشاند که می‌توان از سایر توزیع‌ها به‌جز توزیع نرمال برای متغیر پاسخ استفاده کرد. بنابراین نیازمند روش‌های فراگیری در تحلیل داده‌ها هستیم. در طول دهه‌های اخیر، پیشرفت‌های چشمگیری در تعریف و به‌کارگیری مدل‌ها اتفاق افتاده است. یکی از این پیشرفت‌ها در طول تاریخ رشته آمار، تشخیص این بوده است که می‌توان بسیاری از ویژگی‌های خوب توزیع نرمال را در خانواده توزیع‌های نمایی^۹ نیز یافت. در زیر به معرفی این خانواده و برخی مدل‌های عمومی تأثیرگذار در آن برای تحلیل داده‌های رسته‌ای خواهیم پرداخت.

۱.۴.۱ خانواده توزیع‌های نمایی

متغیر تصادفی Y با مشاهدات $\{y_1, \dots, y_n\}$ را که توزیع احتمال آن به پارامتر θ وابسته است، در نظر بگیرید. این توزیع متعلق به خانواده توزیع‌های نمایی است، اگر تابع (چگالی) احتمال آن به صورت

$$p(y_i|\theta_i) = a(\theta_i)b(y_i) \exp(s(y_i)Q(\theta_i))$$

باشد، که در آن $a(\cdot)$ ، $b(\cdot)$ ، $s(\cdot)$ و $Q(\cdot)$ توابع معلومی هستند. در این توزیع، $Q(\theta_i)$ پارامتر طبیعی نامیده می‌شود. شکل دیگر توزیع به صورت

$$P(Y_i = y_i, \theta_i) = \exp(s(y_i)Q(\theta_i) + c(\theta_i) + d(y_i)) \quad (1.1)$$

است که برای $s(y_i) = y_i$ خانواده را کانونی^{۱۰} می‌نامند (اگرستی، ۲۰۰۲). بسیاری از خانواده توزیع‌های معروف مانند پواسن، نرمال و دوجمله‌ای در این خانواده قرار می‌گیرند. بعضی از آن‌ها را در جدول ۱.۱ می‌توان مشاهده کرد.

^۹Exponential family of distributions

^{۱۰}Canonical

جدول ۱.۱: توزیع‌های پواسن، دوجمله‌ای و نرمال به‌عنوان مثالی از خانواده توزیع‌های نمایی در حالت تک متغیره

توزیع	پارامتر طبیعی	c	d
پواسن	$\log(\theta)$	$-\theta$	$-\log(y!)$
نرمال	$\frac{\mu}{\sigma^2}$	$-\frac{\mu^2}{2\sigma^2} - \frac{1}{2}\log(2\pi\sigma^2)$	$-\frac{y^2}{2\sigma^2}$
دوجمله‌ای	$\log(\frac{\pi}{1-\pi})$	$n\log(1-\pi)$	$\log(\frac{y}{n})$

۲.۴.۱ مدل‌های خطی تعمیم‌یافته

مدل‌های خطی تعمیم‌یافته^{۱۱} شامل دامنه گسترده‌ای از مدل‌های آماری هستند که می‌توان آن‌ها را به مدل‌هایی با متغیرهای پاسخ گسسته تعمیم داد و اغلب در مواردی به کار می‌روند که توزیع متغیر پاسخ نرمال نیست. در این مدل‌ها با پذیرفتن استقلال مشاهدات، با استفاده از یک تابع پیوند بین میانگین مشاهدات و متغیرهای تبیینی ارتباط برقرار می‌شود. همچنین توزیع پاسخ (توزیع خطا) عضوی از خانواده توزیع‌های نمایی و تابع پیوند یک‌نوا و مشتق‌پذیر است. فرض کنید $\{y_i, i = 1, \dots, n\}$ مشاهدات مستقل از یک خانواده از توزیع‌های نمایی باشند که دارای تابع جرم احتمال (۱.۱) هستند. در یک مدل خطی تعمیم‌یافته، مقادیر پارامتر $\theta_i, i = 1, \dots, n$ را به‌صورت تابعی از متغیرهای تبیینی در نظر می‌گیرند. بین پارامتر θ_i و میانگین (شرطی) پاسخ، μ_i ، در GLM رابطه‌ای وجود دارد که توسط تابعی به اسم پیوند تعیین می‌شود. اگر $\{Y_i, i = 1, \dots, n\}$ یک نمونه تصادفی از متغیر پاسخ باشد، آن‌گاه

$$g(\mu_i) = \eta_i = \beta_0 + \sum_{k=1}^l \beta_k x_{ik}, \quad i = 1, \dots, n$$

که در آن β_0 پارامتر عرض از مبدا، β_k ضرایب رگرسیونی و x_k متغیرهای تبیینی هستند و $g(\cdot)$ تابع پیوند است. تابع پیوندی که میانگین را به پارامتر طبیعی متصل کند، پیوند کانونی نامیده می‌شود. در این حالت $g(\mu_i) = Q(\theta_i)$. ترکیب خطی متغیرهای تبیینی، یعنی η ، پیشگوی خطی نامیده می‌شود.

به‌طور خلاصه، مدل‌های خطی تعمیم‌یافته دارای سه مولفه است:

۱. توزیع خطا (یا همان توزیع پاسخ) که عضوی از خانواده توزیع‌های نمایی است و توزیع متغیر پاسخ Y_i را مشخص می‌کند.

۲. مولفه سیستماتیک که متغیرهای تبیینی را به‌صورت ترکیب خطی $\eta_i = x_i' \beta$ در نظر می‌گیرد، که در آن $x_i = (1, x_{i1}, \dots, x_{il})'$ بردار متغیرهای تبیینی برای مشاهده i ام و $\beta = (\beta_0, \beta_1, \dots, \beta_l)'$ بردار ضرایب رگرسیونی است. اغلب مولفه سیستماتیک را پیشگوی خطی می‌گویند.

^{۱۱}Generalized linear models (GLM)

۳. تابع پیوند یکنوا که بین میانگین شرطی توزیع پاسخ و پیشگوی خطی به صورت

$$g(\mu_i) = \eta_i = \mathbf{x}_i' \boldsymbol{\beta},$$

پیوند برقرار می‌کند.

مثال ۱.۴.۱. مدل خطی نرمال: معروفترین مدل خطی تعمیم‌یافته، همان مدل خطی نرمال به صورت

$$Y_i \sim N(\mu_i, \sigma^2), \quad E(Y_i | \mathbf{x}_i) = \mu_i = \mathbf{x}_i' \boldsymbol{\beta},$$

است. در این جا تابع پیوند، تابع همانی $g(\mu_i) = \mu_i$ است. این مدل معمولاً به شکل

$$Y_i = \mathbf{x}_i' \boldsymbol{\beta} + \epsilon_i,$$

نوشته می‌شود، که در آن ϵ_i جمله خطا با توزیع $N(0, \sigma^2)$ است و فرض می‌شود ϵ_i ها مستقل و هم‌توزیع هستند.

مثال ۲.۴.۱. نرخ مرگ و میر: در یک جامعه بزرگ، احتمال پیدا کردن فردی که در یک زمان خاص در حال مرگ باشد، کوچک است. با فرض این که مرگ و میر حاصل از بیماری‌های غیر عفونی مستقل هستند، می‌توان تعداد مرگ و میر، Y ، را با یک توزیع پواسن $Y \sim \text{Pois}(\lambda)$ مدل‌بندی کرد، که در آن $\lambda = E(Y)$ متوسط تعداد مرگ در یک زمان خاص مانند یک سال است. این متوسط می‌تواند به عوامل متعددی مانند اندازه جامعه، سن و جنسیت افراد بستگی داشته باشد. برای وارد کردن اطلاعات متغیرهای تبیینی $\mathbf{x}$ در مدل، می‌توان از مدل

$$\log(\lambda) = \log(n) + \mathbf{x}' \boldsymbol{\beta},$$

استفاده کرد، که در آن n اندازه جامعه و $\mathbf{x}' \boldsymbol{\beta}$ نرخ مرگ به ازای کل جمعیت با اندازه n در هر سال است.

تا این جا دریافتیم مدل خطی تعمیم‌یافته، تعمیم مدل خطی است که در آن مشاهدات می‌توانند گسسته یا رسته‌ای باشند (مک کلاح و نلدر، ۱۹۸۹). در واقع GLM برای یکی کردن مدل‌های آماری مختلف به کار می‌رود تا نتایج جامعی برای رده بزرگی از متغیرهای پاسخ حاصل شود. در مدل‌های خطی، میانگین تابع خطی از پارامترهای رگرسیونی است و واریانس مربوط به مشاهدات ثابت است، در حالی که در مدل خطی تعمیم‌یافته، میانگین به کمک یک تابع پیوند تابعی از متغیرهای تبیینی است. در این صورت، بسیاری از مدل‌های آماری معروف که در آن بین متغیر پاسخ و سایر متغیرهای تبیینی رابطه ناخطی وجود دارد را می‌توان به کمک این تعمیم به کار گرفت. برای تحلیل داده‌های رسته‌ای اغلب از توابع پیوند لوجیت^{۱۲} و پروبیت استفاده می‌کنند که در ادامه مدل‌های مرتبط را معرفی می‌کنیم.

^{۱۲}Logit

تابع پیوند لوجیت

اولین بار بارتلت (۱۹۳۷)، شکل عمومی

$$\log\left(\frac{\pi}{1-\pi}\right) \quad (2.1)$$

را در مدل رگرسیون مطرح کرد. در کتاب جدول‌های آماری که در سال ۱۹۳۸ منتشر شد، فیشر این تبدیل احتمال را برای پارامترهای دوجمله‌ای در تحلیل داده‌های دوجمله‌ای پیشنهاد کرد. در سال ۱۹۴۴، فیزیکدان و آماردان جوزف برکسون اصطلاح لوجیت را برای این تبدیل پیشنهاد داد. مدل لوجیت به‌طور گسترده در علوم زیستی و علوم اجتماعی به‌کار می‌رود. این مدل به‌ویژه برای مطالعات جمعیت‌شناسی برای ارزیابی تأثیر متغیرهای تبیینی بر مخاطره نسبی برآمدها مثل باروری و مرگ و میر کاربرد فراوانی دارد. در ریاضیات و آمار، تابع لوجیت عدد حقیقی $1 > \pi > 0$ به‌صورت

$$\text{logit}(\pi) = \log\left(\frac{\pi}{1-\pi}\right) = \log \pi - \log(1-\pi)$$

تعریف می‌شود. در مدل‌های با توزیع پاسخ دودویی و چندجمله‌ای، استفاده از رگرسیون میانگین برابر با مدل‌بندی روی احتمال است. دامنه رگرسیون میانگین اعداد حقیقی است، در حالی که احتمال مقداری بین صفر و یک است. بنابراین، به کمک مدل خطی تعمیم‌یافته از تابع پیوند لوجیت استفاده می‌کنیم تا بتوانیم دامنه تغییرات میانگین را حفظ کنیم.

مدل لوجیت برای داده‌های دودویی

در تعریف این مدل از منابعی همچون اگوستی (۲۰۱۳)، کانگدان و همکاران (۲۰۰۵) و اگوستی (۲۰۱۹) استفاده کردیم. برای تعریف مدل یک نمونه n تایی را در نظر بگیرید. فرض‌های زیر را برای هر یک از اعضای نمونه می‌پذیریم:

۱. بردار $\mathbf{x}_i = (1, x_{i1}, \dots, x_{il})'$ ، بردار متغیرهای تبیینی i امین عضو از نمونه است.
۲. متغیر Y_i یک پاسخ دودویی برای i امین عضو از نمونه است که مقادیر $\{0, 1\}$ را اختیار می‌کند و از توزیع برنولی پیروی می‌کند. بنابراین

$$Y_i | \mathbf{x}_i \sim \text{Ber}(\pi_i), \quad E(Y_i | \mathbf{x}_i) = P(Y_i = 1 | \mathbf{x}_i) = \pi_i.$$

ساده‌ترین نوع تابعی که بتوان به‌کار گرفت، تابع خطی

$$\pi(\mathbf{x}_i) = \beta_0 + x_{i1}\beta_1 + \dots + x_{il}\beta_l = \mathbf{x}_i' \boldsymbol{\beta},$$

است. اما متأسفانه این مدل ممکن است نتایجی که در آن مقادیر متغیرهای تبیینی و ضرایب رگرسیونی بستگی دارد را به دنبال داشته باشد که در آن π کمتر از صفر یا بیشتر از ۱ شود.

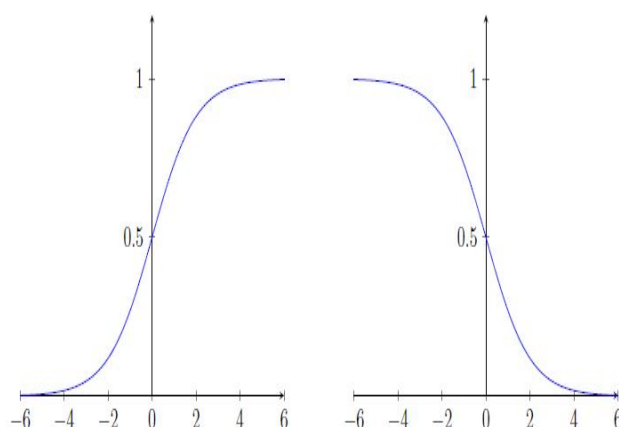
خوشبختانه، توابع پیوند ناخطی متعددی در دسترس هستند که π_i را بین صفر و یک نگه می‌دارند. متداول‌ترین آن‌ها، تابع پیوند لجیت (۲.۱) است. در این صورت

$$\pi(x_i) = \frac{\exp(x'_i\beta)}{1 + \exp(x'_i\beta)}.$$

تفسیر ضرایب رگرسیونی در این مدل بر حسب نسبت بخت‌ها به ازای یک واحد افزایش در متغیر تبیینی، به شرط ثابت بودن سایر متغیرهای تبیینی، است. برای درک این مساله می‌توان نوشت

$$\log \left(\frac{\pi(x_i)}{1 - \pi(x_i)} \right) = x'_i\beta.$$

ضریب مثبت یک متغیر تبیینی باعث افزایش احتمال پیروزی می‌شود. شکل ۱.۱ حالتی را نشان می‌دهد که تنها یک متغیر تبیینی در مدل لجیت حضور دارد. مشاهده می‌کنید که برای مقادیر مثبت ضریب رگرسیونی، احتمال پیروزی افزایشی و برای مقادیر منفی، کاهش می‌یابد. جهت نمودار به علامت x بستگی دارد.



شکل ۱.۱: منحنی رگرسیون احتمال پیروزی به ازای یک متغیر تبیینی، در حالتی که ضریب رگرسیونی مثبت (سمت چپ) یا منفی (سمت راست) است.

مدل لجیت در آزمایش‌های زیستی به‌ویژه برای تعیین میزان مصرف دوز دارو کاربرد گسترده‌ای پیدا کرده است. مدل‌های با پاسخ دوجمله‌ای عامل اصلی انگیزه برای تحلیل داده‌های مربوط به آزمایش‌هایی است که مقدار مختلفی از یک دارو یا مواد آزمایشی ترکیب شده را برای دسته‌ای از موارد مورد آزمایش به کار می‌گیرند. برای مثال، فرض کنید در آزمایشی برای تعیین دوز حشره‌کش، تعدادی از حشره‌کش‌ها را برای تعیین سطح دوز مصرفی در سطح معینی می‌آزمایند. در روزهای ابتدایی هیچکدام یا برخی از حشره‌ها از بین نمی‌روند. اما در روزهای بعد ممکن است نابود شوند که حاکی از تاثیر دوز مصرفی بر روی آستانه تحمل حشره دارد.

مدل لجیت برای داده‌های چندجمله‌ای

در تعریف این مدل از اگوستی (۲۰۱۳)، دو و همکاران (۲۰۰۴)، کانگدان (۲۰۰۵) و سید (۲۰۲۰) استفاده کردیم. داده‌های اسمی با J رشته را در نظر بگیرید. در این حالت، مدل لجیت لگاریتم بخت، یعنی $\log(\frac{\pi}{1-\pi})$ ، را برای هر جفت از رده‌ها در نظر می‌گیرد. برای متغیرهای تبیینی $\mathbf{x}' = (x_1, \dots, x_l)$ فرض کنید

$$\pi_{ij} = \pi_j(\mathbf{x}_i) = P(Y_i = j | \mathbf{x}_i), \quad j = 1, \dots, J; \quad \sum_{j=1}^J \pi_{ij} = 1.$$

در حالت کلی برای داده‌های با پاسخ چند رده‌ای، لگاریتم بخت را به صورت مقایسه زوجی بین رده‌ها به کار می‌گیرند. برای مثال، $\frac{\pi_{ij}}{\pi_{iJ}}$ بخت موفقیت بین دو رده j و J است. در این حالت به رده J پایه می‌گویند که انتخاب آن دلخواه است. بنابراین برای هر $j = 1, \dots, J$ که $j \neq J$ و $\beta_j = (\beta_{1j}, \beta_{2j}, \dots, \beta_{lj})'$ داریم

$$\begin{aligned} \log\left(\frac{\pi_{ij}(\mathbf{x}_i)}{\pi_{iJ}(\mathbf{x}_i)}\right) &= \log\left(\frac{\pi_{ij}(\mathbf{x}_i)}{1 - \sum_{j=1}^{J-1} \pi_{ij}(\mathbf{x}_i)}\right) = \beta_{0j} + \beta_{1j}x_{i1} + \dots + \beta_{lj}x_{il} \quad (3.1) \\ &= \beta_{0j} + \sum_{k=1}^l x_{ik}\beta_{kj} = \beta_{0j} + \mathbf{x}_i'\beta_j = \eta_{ij} \end{aligned}$$

و همچنین

- برای هر i ، مشاهده مربوط به j امین رشته از متغیر تصادفی Y_{ij} است.
- احتمال مشاهده j امین رشته از متغیر پاسخ برای i امین مورد، π_{ij} است و
- ضریب رگرسیونی k امین متغیر تبیینی است.

با بازنویسی معادله (۳.۱) به کمک ویژگی‌های تابع لگاریتم، می‌توان نوشت

$$\pi_{ij} = \pi_{iJ} \exp(\eta_{ij}).$$

با توجه به این که $\sum_{j=1}^J \pi_{ij} = 1$ ، پس

$$1 = \pi_{iJ} \exp(\beta_{01} + \mathbf{x}_i'\beta_1) + \pi_{iJ} \exp(\beta_{02} + \mathbf{x}_i'\beta_2) + \dots + \pi_{iJ}.$$

بنابراین

$$\pi_{iJ} = \frac{1}{1 + \sum_{j=1}^{J-1} \exp(\beta_{0j} + \mathbf{x}_i'\beta_j)}.$$

از این رو، عبارت عمومی برای π_{ij} ، به شکل

$$\pi_{ij} = \frac{\exp(\beta_{0j} + \mathbf{x}_i'\beta_j)}{1 + \sum_{j=1}^{J-1} \exp(\beta_{0j} + \mathbf{x}_i'\beta_j)} = \frac{\eta_{ij}}{1 + \sum_{j=1}^{J-1} \eta_{ij}}$$

به دست می‌آید. برای تشریح مدل به مثال زیر برگرفته از اگوستی (۲۰۱۳) توجه کنید.

مثال ۳.۴.۱. پژوهشگران می‌خواهند عواملی را که بر انتخاب غذا توسط سوسمار تأثیر می‌گذارند، بیابند. برای این منظور ۲۱۹ سوسمار از ۴ دریاچه در ایالت فلوریدا را انتخاب کردند. متغیر پاسخ اسمی (یعنی نوع غذای دلخواه) را بر اساس بیشترین حجم غذای اولیه‌ای که از شکم سوسمار می‌یابند، تعیین می‌کنند. با این توصیف، متغیر پاسخ دارای پنج رده ماهی (F)، پرنده‌ها (B)، بی‌مهرگان (I)، خزنده‌ها (R) و سایر (O) است. سوسمارها در دو اندازه مختلف از سه دریاچه مختلف انتخاب شدند. بنابراین اندازه سوسمار و نوع دریاچه را به‌عنوان متغیرهای تبیینی می‌توان در نظر گرفت. متغیرهای تبیینی s، h، o و t به ترتیب نشان‌دهنده اندازه سوسمار، دریاچه اول، دریاچه دوم و دریاچه سوم هستند که همه متغیرهای نشانگر دوسطحی هستند (اندازه سوسمارها به دو رده اندازه بزرگ و کوچک دسته‌بندی شده‌اند). با در نظر گرفتن ماهی به‌عنوان رده پایه، نتایج برآوردهای ماکسیمم درست‌نمایی پارامترهای مدل لجیت چندجمله‌ای در جدول ۲.۱ گزارش شده‌اند. مقادیر خطای معیار برآورد اثرات متغیرهای تبیینی داخل پرنتر جلوی هر برآورد ضریب نوشته شده‌اند. برای مثال، معادله پیشگو برای مقایسه

جدول ۲.۱: پارامترهای برآورده شده مدل لجیت چندجمله‌ای در مثال سوسمارها

مدل لجیت	عرض از مبدأ	اثر اندازه	اثر دریاچه اول	اثر دریاچه دوم	اثر دریاچه سوم
$\text{logit}(\frac{\pi_I}{\pi_F})$	-۱/۵۵	۱/۴۶ (۰/۴۰)	-۱/۶۶ (۰/۶۱)	۰/۹۴ (۰/۴۷)	۱/۱۲ (۰/۴۹)
$\text{logit}(\frac{\pi_R}{\pi_F})$	-۳/۳۱	-۰/۳۵ (۰/۵۸)	۱/۲۴ (۱/۱۹)	۲/۴۶ (۱/۱۲)	۲/۹۴ (۱/۱۲)
$\text{logit}(\frac{\pi_B}{\pi_F})$	-۲/۰۹	-۰/۶۳ (۰/۶۴)	۰/۷۰ (۰/۷۸)	-۰/۶۵ (۱/۲۰)	۱/۰۹ (۰/۸۴)
$\text{logit}(\frac{\pi_O}{\pi_F})$	-۱/۹۰	۰/۳۳ (۰/۴۵)	۰/۸۳ (۰/۵۶)	۰/۰۱ (۰/۷۸)	۱/۵۲ (۰/۶۲)

خوراک بی‌مهرگان در مقابل با ماهی‌ها به‌صورت زیر است:

$$\text{logit}\left(\frac{\hat{\pi}_I}{\hat{\pi}_F}\right) = -1/55 + 1/46s - 1/66h + 0/94o + 1/12t.$$

در این مدل، به ازای دریاچه مشخص، احتمال انتخاب غذای بی‌مهرگان برای سوسمارهای با اندازه کوچک $\exp(1/46) = 4/3$ برابر زمانی است که اندازه سوسمار بزرگ باشد. همچنین، احتمال انتخاب خوراک بی‌مهرگان در دریاچه‌های دوم و سوم بیشتر از دریاچه اول است.

مدل لجیت برای داده‌های ترتیبی

در بیان این مدل از مک کلاک (۱۹۸۰) استفاده کردیم. در تحلیل داده‌های ترتیبی به کمک تابع پیوند لجیت، از شکل تجمعی آن بهره می‌گیرند. فرض کنید برای متغیر تصادفی رسته‌ای Y ، تعداد رسته‌ها باشد. اگر نمونه‌های $n_1, \dots, n_J$ را فراوانی در هر رسته در نظر بگیریم، تعداد کل مشاهدات $n = \sum_j n_j$ است. برای یک مشاهده که به‌طور تصادفی از جامعه انتخاب شده است، فرض کنید π_j احتمال پاسخ در رسته j باشد. احتمال تجمعی رده j به‌صورت

$$F_j = P(Y \leq j) = \pi_1 + \dots + \pi_j, \quad j = 1, 2, \dots, J$$

است که ترتیب رسته‌ها را به صورت

$$0 < F_1 < F_2 < \dots < F_J = 1$$

نشان می‌دهد.

برای J رسته با احتمال‌های $\pi_1, \dots, \pi_J$ ، مدل لوجیت تجمعی^{۱۳} به صورت

$$\begin{aligned} \text{logit}\{P(Y \leq j)\} &= \log \frac{P(Y \leq j)}{1 - P(Y \leq j)} \\ &= \log \frac{\pi_1 + \dots + \pi_j}{\pi_{j+1} + \dots + \pi_J}, \quad j = 1, \dots, J-1 \end{aligned}$$

تعریف می‌شود. این مدل باعث می‌شود تا پاسخ‌ها به دو قسمت $Y \leq j$ و $Y > j$ تفکیک شوند که خود، نوعی داده ترتیبی دودویی به شمار می‌آید. هر لوجیت تجمعی از اطلاعات همه J رسته استفاده می‌کند.

مدل لوجیت تجمعی با متغیر تبیینی

فرض کنید Y_i متغیر پاسخ رسته‌ای با J رسته و x_i بردار متغیرهای تبیینی باشد. مدل لوجیت تجمعی در این حالت، هم‌زمان از لوجیت تجمعی $(J-1)$ رسته استفاده می‌کند که به صورت

$$\text{logit}\{P(Y_i \leq j)\} = \beta_{0j} + x_i' \beta, \quad j = 1, \dots, J$$

تعریف می‌شود. این مدل یک احتمال شرطی به شرط متغیرهای تبیینی است، بنابراین باید در واقع $P(Y_i \leq j)$ را با $P(Y_i \leq j | x_i)$ جایگزین کنیم. یک مدل معادل برای لوجیت تجمعی به صورت

$$P(Y_i \leq j) = \frac{\exp(\beta_{0j} + x_i' \beta)}{1 + \exp(\beta_{0j} + x_i' \beta)}$$

است. در این مدل‌ها بر خلاف داده‌های رسته‌ای اسمی، پارامترهای رگرسیونی برای همه رسته‌ها یکسان است.

مثال ۴.۴.۱. در تحلیل اختلالات روانی افراد در یک آسایشگاه روانی، شدت اختلالات روانی افراد یک متغیر پاسخ ترتیبی است. رده‌های پاسخ عبارتند از خوب، با علائم خفیف، با علائم متوسط و با اختلال روانی جدی. متغیرهای تبیینی نیز شامل سابقه‌ای از وقایع زندگی افراد (x_1) و وضعیت اجتماعی-اقتصادی (x_2) هستند. مدل لوجیت تجمعی به صورت

$$\text{logit}\{P(Y_i \leq j)\} = \beta_{0j} + \beta_1 x_1 + \beta_2 x_2, \quad j = 1, 2, 3, 4,$$

است. برآورد ماکسیمم درست‌نمایی ضرایب رگرسیونی β_1 و β_2 به ترتیب برابر با -0.319 و $1/111$ است. برآوردها نشان می‌دهد که داشتن وضعیت اجتماعی-اقتصادی بهتر (یعنی

^{۱۳} Cumulative

مشاهده بزرگتر برای این متغیر) احتمال تجمعی را افزایش می‌دهد و این به آن معنی است که احتمال داشتن شدت اختلال بالاتر کم خواهد شد. نتیجه برای متغیر تبیینی دوم، با توجه به ضریب منفی آن، برعکس است. داده‌های این مثال در <http://www.stat.ualberta.ca> قابل دسترسی هستند.

تابع پیوند پروبیت

نوع دیگری از تابع پیوند برای تحلیل داده‌های رسته‌ای، تابع پیوند پروبیت^{۱۴} است که مدل نتیجه‌شده یکی از مدل‌های معروف در تحلیل داده‌های رسته‌ای محسوب می‌شود. این تابع پیوند، اولین بار توسط فردی به نام بلیس در سال ۱۹۳۴ معرفی شد.

مدل پروبیت برای داده‌های دودویی

در این حالت نیز مدل ناخطی در π به یک مدل خطی از π نسبت به متغیرهای تبیینی توسط یک تابع یکنوا تغییر شکل می‌دهد. احتمال i امین مشاهده یعنی π_i ، با استفاده از توزیع نرمال استاندارد به صورت

$$\pi_i = \int_{-\infty}^{\eta_i} \frac{1}{\sqrt{2\pi}} \exp\left(-\frac{1}{2}u^2\right) du = \Phi(\eta_i),$$

است که از آن داریم

$$\eta_i = \Phi^{-1}(\pi_i).$$

بنابراین مدل پروبیت را می‌توان به صورت

$$\Phi^{-1}(\pi_i) = \eta_i = \beta_0 + \sum_{k=1}^l x_{ik}\beta_k,$$

یا

$$\pi_i = \Phi\left(\beta_0 + \sum_{k=1}^l x_{ik}\beta_k\right),$$

نوشت.

مدل پروبیت تجمعی برای داده‌های رسته‌ای ترتیبی

مدل پروبیت تجمعی به صورت

$$\Phi^{-1}(P(Y_i \leq j)) = \beta_{0j} + x'_i\beta, \quad j = 1, \dots, J-1,$$

است که در علوم مختلف مدل پروبیت ترتیبی نامیده می‌شود. همانند مدل لجیت، اثر β برای همه رسته‌ها یکسان است. مدل پروبیت تجمعی، احتمالات تجمعی را به صورت

$$P(Y_i \leq j) = \Phi(\beta_{0j} + x'_i\beta), \quad j = 1, \dots, J-1,$$

^{۱۴}Probit

توصیف می‌کند.

۵.۱ ساختارهای وابستگی

تا به اینجا دانستیم که می‌توان به کمک رده مدل‌های خطی تعمیم‌یافته، بسیاری از داده‌های رسته‌ای را با فرض این که احتمال وقوع یک رسته تابعی از متغیرهای تبیینی باشد، مدل‌بندی و تحلیل کرد. اگرچه پیشرفت بزرگی به‌نظر می‌رسد، اما گاهی اوقات با داده‌هایی روبرو هستیم که در بین آن‌ها ساختار همبستگی وجود دارد. در این حالت، باید به دنبال مدلی باشیم که همبستگی موجود در داده‌ها را برای به‌دست آوردن کارایی بالاتر در تحلیل لحاظ کند. یکی از مدل‌های مناسب، مدل آمیخته خطی تعمیم‌یافته^{۱۵} (فارمیر و توتز، ۲۰۰۱) نام دارد.

۱.۵.۱ مدل‌های آمیخته خطی تعمیم‌یافته

در حالتی که بین مشاهدات همبستگی وجود دارد، در مدل خطی تعمیم‌یافته، پذیره استقلال مشاهدات به استقلال شرطی تعدیل و همبستگی بین آن‌ها با اضافه کردن اثرات تصادفی^{۱۶} از طریق متغیرهای پنهان^{۱۷} به مدل در نظر گرفته می‌شود. در این صورت مدل خطی تعمیم‌یافته به مدل آمیخته خطی تعمیم‌یافته تبدیل می‌شود و پیشگوی خطی در آن، در قالب ماتریسی، به‌صورت

$$\eta = X\beta + Zu,$$

بیان می‌شود، که در آن $\eta' = (\eta_1, \dots, \eta_n)$ بردار $l+1$ بعدی از اثرات ثابت، u بردار q بعدی از اثرات تصادفی، X ماتریس طرح $n \times (l+1)$ بعدی و Z ماتریس طرح $n \times q$ بعدی است. معمولاً در این رده از مدل‌ها، فرض می‌شود بردار u دارای توزیع گاوسی است و وابستگی مشاهدات پاسخ از طریق همین بردار تصادفی وارد مدل می‌شود. در حالت خاصی که $q = 1$ و Z یک بردار n بعدی از ۱‌هاست، مدل حاصل به مدل با عرض از مبدا تصادفی^{۱۸} معروف است.

۲.۵.۱ مدل آمیخته جمعی-ساختاری تعمیم‌یافته

با پیشرفت‌های مختلفی که در علوم رخ داده است، داده‌هایی در طبیعت مشاهده می‌کنیم که همواره در بین آن‌ها رد پای اثرات مختلف را می‌توان جستجو کرد. برای مثال، در تحلیل داده‌ها ممکن است برخی متغیرهای تبیینی به‌صورت ناخطی روی متغیر پاسخ تاثیر داشته باشند. ممکن است اثرات متقابل در بین برخی از متغیرها وجود داشته باشد که تغییرات مربوط

^{۱۵}Generalized linear mixed model (GLMMs)

^{۱۶}Random effects

^{۱۷}Latent variables

^{۱۸}Random intercept model

به آن‌ها نیز هموار نیست. همچنین در بعضی از داده‌ها اثر مکان رخداد پاسخ یا نقطه‌ای که داده متعلق به آن است را می‌توان یافت. افزون بر این، از اثرات مربوط به قرار گرفتن مشاهدات در یک بازه زمانی خاص نمی‌توان غافل شد. به‌طور مثال، اگر داده‌ها مربوط به مقایسه پیشرفت یا مبتلا شدن افراد به یک بیماری خاص در یک بازه زمانی مشخص و در ناحیه‌های مختلف یک کشور باشد، آن‌گاه مدل‌بندی مربوط به تأثیر اقلیم و نواحی محل زندگی در طول زمان مبتلا به این بیماری، گریزنپذیر است. در این صورت باید به دنبال مدلی باشیم تا از همه این اطلاعات بهره ببرد.

هستی و تیبشرانی (۱۹۹۰) رده مدل‌های جمعی تعمیم‌یافته^{۱۹} را معرفی کردند که می‌توان در آن اثر ناخطی متغیرهای تبیینی را وارد مدل کرد. با برقراری پذیره‌های مدل خطی تعمیم‌یافته، در این رده از مدل‌ها، پیشگوی خطی به یک پیشگوی جمعی به‌صورت

$$g(\mu) = \eta = \beta_0 + \sum_{j=1}^l f_j(x_j),$$

تبدیل می‌شود که در آن $f_j(\cdot)$ اثرات ناخطی (ناپارامتری) متغیرهای تبیینی هستند که معمولاً شرایط همواری مشخصی برای آن‌ها در نظر گرفته می‌شود. تعمیمی دیگر از مدل GAM مدلی است که شامل هر دو اثرات خطی و ناخطی به شکل

$$g(\mu) = \eta = \beta_0 + \mathbf{x}'_i \boldsymbol{\beta} + \sum_{c=1}^p f_c(z_c),$$

است، که در آن $\mathbf{x}$ دارای اثرات خطی و z_c ها دارای اثرات ناخطی هستند. این نوع مدل، در واقع، یک مدل نیمه‌پارامتری است.

تعمیمی دیگر بر مدل نیمه‌پارامتری، تحت عنوان مدل آمیخته جمعی – ساختاری تعمیم‌یافته^{۲۰}، به پیشگوی جمعی اثرات تصادفی را نیز اضافه می‌کند که برای پاسخ‌های با ساختار وابستگی می‌تواند مناسب و به اندازه کافی منعطف باشد. در این حالت

$$g(\mu) = \eta = \beta_0 + \mathbf{x}'_i \boldsymbol{\beta} + \sum_{c=1}^p f_c(z_c) + f(u), \quad (4.1)$$

که در آن $f(u)$ اثر تصادفی با معمولاً توزیع نرمال است. معمولاً یک حالت کلی‌تر برای پیشگوی (۴.۱) مطرح است که به آن مدل رگرسیون جمعی – ساختاری^{۲۱} می‌گویند (فاهرمیر و توتز، ۲۰۰۱). در این نوع از مدل‌بندی، پیشگوی خطی با یک پیشگوی کامل‌تر جمعی – ساختاری به صورت

$$\eta = \beta_0 + \sum_{r=1}^{n_f} f(u_r) + \sum_{k=1}^l x_k \beta_k,$$

^{۱۹}Generalized additive models (GAMs)

^{۲۰}Generalized additive mixed models (GAMMs)

^{۲۱}Structured additive regression model (STAR)

جایگزین می‌شود. به کمک توابع $f(\cdot)$ می‌توان اثرات ساختاری مختلف را وارد مدل کرد. این اثرات می‌توانند شامل

- اثرات ناخطی توابع پیوسته
- رویه‌های هموار دوبعدی (یا چندبعدی)
- اثر فضایی
- اثر فضایی-زمانی
- متغیرهای با اثرات متغیر زمانی یا فضایی

شوند. در این مدل، متغیرهای u می‌توانند متغیرهای اندیس‌گذار زمان، مکان، زمان-مکان یا مشاهدات یک متغیر تبیینی پیوسته باشند. با معرفی مدل‌های بالا، دریافتیم که برای لحاظ کردن اثرات فضایی و فضایی-زمانی (که مد نظر این رساله است) در تحلیل داده‌های رسته‌ای، باید از چه رویکردهایی در مدل‌بندی استفاده کنیم. در ادامه، به‌طور خلاصه به تعریف این اثرها خواهیم پرداخت.

۶.۱ ساختارهای وابستگی خاص

در فرآیند تحلیل داده‌ها، ممکن است با داده‌هایی روبرو شویم که از نظر مکانی و زمانی به یکدیگر وابسته هستند. به‌ویژه این که در دهه‌های اخیر، به‌خاطر پیشرفت در ابزارهایی مانند ماهواره‌ها و GPS، جمع‌آوری داده‌های واقعی زمانی ممکن شده است. در واقع، امروزه در شاخه‌های علمی مختلفی مانند محیط زیست، اقلیم‌شناسی، علوم اجتماعی و علوم پزشکی، داده‌ها همراه اطلاعات فضا و زمان رخداد آن‌ها در اختیار محققان قرار دارند. به‌عنوان مثال، فرض کنید میزان وقوع یک بیماری خاص مثلاً سرطان ریه در تمام کشور یا نواحی بزرگی از آن در سال‌های مختلف، هدف مطالعه باشد. به نظر شما چه نوع مدلی این امکان را به ما می‌دهد تا جای ممکن از تمام اطلاعات موجود در داده‌ها به خوبی استفاده کنیم. برای این کار، ابتدا باید الگوی جغرافیایی بیماری شناسایی و تعیین شود، به‌طوری که ناحیه‌های نزدیک به هم الگوهای مشابهی از شاخص‌های جغرافیایی را نمایش دهند. این وظیفه توسط مدل‌های فضایی^{۲۲} انجام می‌شود. به عبارتی، احتمال وقوع رخداد بیماری در ناحیه‌های همسایه به هم نزدیک‌تر است. همچنین سوال مهم دیگر این می‌تواند باشد که احتمال وقوع در طول زمان چگونه تغییر می‌کند؟ یا این که بخواهیم مسئله را هم از نظر جغرافیایی و هم از نظر زمانی بررسی کنیم که نتیجه به مدل‌های فضایی-زمانی^{۲۳} منتهی می‌شود. امروزه مدل‌های فضایی

^{۲۲} Spatial models

^{۲۳} Spatio-temporal models

و فضایی-زمانی به‌طور وسیعی مورد استفاده قرار می‌گیرند. از همین رو، مقدمه مختصری از برخی مدل‌های مربوط به ساختار وابستگی که امروزه به وفور گسترش یافته‌اند، بیان می‌کنیم. این مدل‌ها در فصل‌های آینده به کار می‌روند.

۱.۶.۱ ساختار وابستگی زمانی

در بسیاری از موقعیت‌های کاربردی با داده‌هایی مواجه می‌شویم که از نظر زمانی به یکدیگر وابسته هستند. رویکردهای مختلفی برای مدل‌بندی ساختار وابستگی زمانی معرفی شده‌اند. به‌عنوان نمونه، ساختار وابستگی زمانی را می‌توان با مدل‌های اتورگرسیو^{۲۴}، میانگین متحرک^{۲۵} یا ترکیب آن‌ها یعنی اتورگرسیو میانگین متحرک^{۲۶} وارد فرآیند تحلیل داده‌ها کرد (باکس و همکاران، ۱۹۹۴). به‌عنوان مثال، فرآیند اتورگرسیو به صورت زیر تعریف می‌شود.

تعریف ۱.۶.۱. دنباله $\{X_t, t \in \mathcal{T}\}$ یک فرآیند اتورگرسیو مرتبه p است، هرگاه بتوان آن را به صورت

$$X_t = \psi_1 X_{t-1} + \psi_2 X_{t-2} + \dots + \psi_p X_{t-p} + \epsilon_t$$

نوشت، که در آن ψ ها پارامترهای مدل، ϵ_t یک فرآیند نوفه سفید (نرمال) و $\mathcal{T}$ مجموعه (گسسته) اندیس‌گذار زمانی است. زمانی که p برابر با یک باشد، فرآیند اتورگرسیو از مرتبه اول است و ساختار وابستگی ماتریس کوواریانس، Σ ، به صورت $\sigma^2 \rho^{|i-j|}$ است، که در آن σ^2 واریانس فرآیند نوفه سفید است. در این حالت (برای سه نقطه زمانی گسسته)

$$\Sigma = \begin{bmatrix} 1 & \rho & \rho^2 & \rho^3 \\ & 1 & \rho & \rho^2 \\ & & 1 & \rho \\ & & & 1 \end{bmatrix}.$$

از سایر مدل‌های زمانی معروف و پرکاربرد، می‌توان به فرآیندهای قدم‌زدن تصادفی مرتبه اول^{۲۷} (RW1) و مرتبه دوم (RW2) اشاره کرد (رو و هلد، ۲۰۰۵).

تعریف ۲.۶.۱. دنباله $\{\nu_t, t \in \mathcal{T}\}$ یک فرآیند قدم‌زدن تصادفی مرتبه اول است، هرگاه رابطه

$$\nu_t | \nu_{t-1} \sim N(\nu_{t-1}, \sigma),$$

برقرار باشد که در آن σ انحراف معیار شرطی مدل است. به‌طور مشابه، اگر

$$\nu_t | \nu_{t-1}, \nu_{t-2} \sim N(2\nu_{t-1} - \nu_{t-2}, \sigma),$$

^{۲۴} Autoregressive (AR)

^{۲۵} Moving-average (MA)

^{۲۶} Autoregressive moving-average (ARMA)

^{۲۷} First order random walk

آن گاه فرآیند را قدم‌زدن تصادفی مرتبه دوم می‌نامند. دقت کنید که مثلاً برای مدل قدم‌زدن تصادفی مرتبه اول $\nu_t | \nu_{t-1}, \dots$ با $\nu_t | \nu_{t-1}$ هم‌توزیع است.

۲.۶.۱ ساختار وابستگی فضایی

در بررسی‌های محیطی، اغلب با داده‌هایی مواجه می‌شویم که از یکدیگر مستقل نیستند و وابستگی آن‌ها ناشی از موقعیت و مکان قرار گرفتن داده‌ها در فضای مورد بررسی است. این گونه داده‌ها، داده‌های فضایی نامیده می‌شوند. داده‌های زمین آماری^{۲۸} و داده‌های شبکه‌ای^{۲۹} از انواع داده‌های فضایی محسوب می‌شوند.

داده‌های زمین آماری

این نوع داده‌ها در موقعیت‌های ثابت و مشخص در ناحیه‌ای پیوسته (چگال) مشاهده می‌شوند. متغیر مورد بررسی ممکن است گسسته یا پیوسته باشد. به عنوان مثال، برای حالت پیوسته می‌توان غلظت مواد معدنی در یک معدن و مقدار ریزش باران ثبت شده در ایستگاه‌های هواشناسی و برای حالت گسسته می‌توان تعداد جانوران دریایی در یک مجموعه از مکان‌های نمونه‌گیری شده در طول یک ساحل را نام برد (محمدزاده، ۱۳۹۱). برای مدل‌بندی داده‌های زمین آماری از میدان تصادفی استفاده می‌کنند که در زیر معرفی شده است.

تعریف ۳.۶.۱. میدان تصادفی، مجموعه‌ای از متغیرهای تصادفی مانند $\{Y(s); s \in D\}$ است که در آن مجموعه اندیس‌گذار D یک زیرمجموعه از فضای اقلیدسی d بعدی، $d \geq 1$ است. برای این میدان تصادفی، میانگین در موقعیت s و کوواریانس در موقعیت‌های s_1 و s_2 به ترتیب به صورت

$$\begin{aligned}\mu(s) &= E(Y(s)), \quad s \in D, \\ C(s_1, s_2) &= \text{Cov}(Y(s_1), Y(s_2)), \quad s_1, s_2 \in D,\end{aligned}$$

تعریف می‌شوند.

اگر میدان تصادفی نرمال^{۳۰} باشد، آن گاه برای هر $n \geq 1$ و برای هر مجموعه از موقعیت‌های $\{s_1, \dots, s_n\}$ بردار $(Y(s_1), \dots, Y(s_n))'$ دارای توزیع نرمال n متغیره با بردار میانگین $\mu = (\mu(s_1), \dots, \mu(s_n))'$ و ماتریس کوواریانس Σ است که مولفه‌های آن توسط تابع کوواریانس $C(\cdot, \cdot)$ تعیین می‌شوند.

^{۲۸}Geostatistical data

^{۲۹}Lattice data

^{۳۰}Gaussian random field (GRF)

ماتریس کوواریانس با ساختار فضایی را می‌توان به کمک توابع کوواریانس فضایی پارامتری مانند مترن^{۳۱}، نمایی، گوسی، کروی^{۳۲}، دایره‌ای و درجه ۳ به‌دست آورد. به‌عنوان مثال، تابع کوواریانس مترن برای موقعیت‌های s_i و s_j به‌صورت

$$\text{Cov}(Y(s_i), Y(s_j)) = \text{Cov}(Y_i, Y_j) = \frac{\sigma^2}{\Gamma(\lambda) 2^{\lambda-1}} (\rho \|s_i - s_j\|)^\lambda \kappa_\lambda(\rho \|s_i - s_j\|),$$

تعریف می‌شود، که در آن، $\|s_i - s_j\|$ فاصله اقلیدسی بین موقعیت‌های s_i و s_j ، σ^2 واریانس کناری میدان مترن، $\kappa_\lambda(\cdot)$ تابع بسل^{۳۳} نوع دوم (آبراموویچ و همکاران، ۱۹۷۲) با مرتبه $\lambda > 0$ و $\Gamma(\cdot)$ تابع گاما است. پارامتر λ نیز درجه همواری فرآیند را اندازه‌گیری می‌کند و معمولاً ثابت نگه داشته می‌شود. همچنین ρ پارامتر مقیاس است. در عمل پارامترهای تابع کوواریانس بر اساس مشاهداتی از میدان و با روش‌های درست‌نمایی ماکسیمم یا بیزی برآورد می‌شوند (کرسی، ۱۹۹۳؛ محمدزاده، ۱۹۹۱).

داده‌های شبکه‌ای

این نوع داده‌ها مربوط به مکان‌های ناحیه‌ای هستند، که ممکن است منظم یا نامنظم باشند. تصاویر ماهواره‌ای از سطح زمین که به تعدادی عنصر تصویر^{۳۴} کوچک به‌صورت شبکه منظم در $\mathbb{R}^2$ تقسیم شده‌اند، مثالی برای داده‌های شبکه‌ای منظم است. تعداد حوادث رانندگی، تعداد اتباع خارجی غیر مجاز ساکن یا نرخ رشد اقتصادی هر استان مثال‌های دیگر از داده‌های شبکه‌ای هستند (محمدزاده، ۱۳۹۱). برای مدل‌بندی ساختار فضایی داده‌های شبکه‌ای می‌توان از مدل اتورگرسیو شرطی^{۳۵}؛ بی‌سیگ، ۱۹۷۴) استفاده کرد. مدل CAR برای بردار تصادفی $(Y_1, \dots, Y_n)'$ که دارای توزیع نرمال باشد، به صورت زیر تعریف می‌شود، که در آن چگالی‌های شرطی تعیین می‌شوند و از روی آن‌ها می‌توان پارامترهای توزیع توام بردار را کاملاً به‌صورت

$$Y_i | \mathbf{Y}_{-i}, \tau, \mathbf{W} \sim N\left(\frac{\sum_{m \sim i} \omega_{im} Y_m}{\sum_{m \sim i} \omega_{im}}, \frac{1}{\tau \sum_{m \sim i} \omega_{im}}\right) \quad (5.1)$$

مشخص کرد، که در آن $\mathbf{Y}_{-i}$ بردار بدون مولفه i ام است، منظور از $i \sim m$ همسایه بودن دو ناحیه i و m است، τ پارامتر دقت مدل است و ماتریس همسایگی $\mathbf{W} = \{\omega_{im}, i \sim m, i > m\}$ دارای مولفه‌های صفر و یک است که یک‌ها مشخص می‌کنند هر ناحیه با کدام نواحی دیگر همسایه است. در این حالت امید شرطی، میانگین اثرهای تصادفی در نواحی همسایه و واریانس شرطی، عکس نسبت همسایه‌ها است.

^{۳۱} Matern

^{۳۲} Spherical

^{۳۳} Bessel

^{۳۴} Pixel

^{۳۵} Conditional autoregressive (CAR)

این مدل با میانگین‌گیری اثرات همسایه‌ها، نوعی همواری را در میدان تصادفی القا می‌کند که برای کاربردهای آمار فضایی شبکه‌ای مطلوب هم هست. مدل (۵.۱) با توجه به این که یک چگالی سره را نتیجه نمی‌دهد، به مدل اتورگرسیو شرطی ذاتی^{۳۶} معروف است (رو و هلد، ۲۰۰۵). نسخه‌های سره برای مدل‌های CAR نیز وجود دارند که استفاده از مدل ICAR که یکی از این نسخه‌ها است، در اغلب کاربردها برتری دارد. برای جزئیات بیشتر می‌توانید به کتاب رو و هلد (۲۰۰۵) مراجعه کنید.

۳.۶.۱ وابستگی فضایی-زمانی

تحلیل داده‌های فضایی-زمانی مستلزم تعیین ساختار همبستگی فضایی-زمانی آن‌ها از طریق تابع کوواریانس است. در دو دهه اخیر بررسی‌های زیادی برای تعیین ساختار همبستگی، مدل‌بندی و تحلیل چنین داده‌هایی انجام شده‌اند (ویکل و کرسی، ۲۰۱۱). ساختار همبستگی این داده‌ها به وسیله کوواریانس فضایی-زمانی تعیین می‌شود. این تابع که نقش به‌سزایی در پیشگویی موقعیت‌های فضایی یا زمانی فاقد مشاهده دارد، به‌طور معمول نامعلوم است و باید بر اساس مشاهدات برآورد شود. ساخت توابع کوواریانس فضایی-زمانی معتبر را می‌توان به دو صورت تفکیک‌پذیر^{۳۷} و تفکیک‌ناپذیر بررسی کرد که معمولاً برای سهولت در برازش مدل معتبر از صورت تفکیک‌پذیر آن استفاده می‌کنند.

تعریف ۴.۶.۱. یک میدان تصادفی فضایی-زمانی، مجموعه‌ای از متغیرهای تصادفی مانند $Y(\cdot, \cdot) = \{Y(s, t); (s, t) \in D \times T\}$ است، که در آن s موقعیت فضایی در مجموعه $D \subseteq \mathbb{R}^d$ ، $d \geq 1$ و t لحظه زمانی در مجموعه $T \in \mathbb{R}^+$ است.

تعریف ۵.۶.۱. میدان تصادفی فضایی-زمانی $Y(\cdot, \cdot)$ تفکیک‌پذیر نامیده می‌شود، اگر توابع کوواریانس فضایی و زمانی C_s و C_r وجود داشته باشند، به‌طوری که

$$\text{Cov}(Y(s, t), Y(s', t')) = C_s(s, s')C_r(t, t'), \quad (s, t), (s', t') \in D \times T.$$

برای یک میدان تصادفی فضایی-زمانی تفکیک‌پذیر، ماتریس کوواریانس Σ برابر حاصلضرب کرونه‌کری ماتریس کوواریانس فضایی $n \times n$ بعدی $\Sigma_s = (\sigma_{ij})$ و ماتریس کوواریانس زمانی $m \times m$ بعدی Σ_t به صورت

$$\Sigma = \Sigma_s \otimes \Sigma_t = \begin{bmatrix} \sigma_{11}\Sigma_t & \cdots & \sigma_{1n}\Sigma_t \\ \vdots & \ddots & \vdots \\ \sigma_{n1}\Sigma_t & \cdots & \sigma_{nn}\Sigma_t \end{bmatrix}$$

^{۳۶}Intrinsic CAR (ICAR)

^{۳۷}Seperable

است. با انتخاب یک مدل معتبر برای Σ_s که در بخش ساختار وابستگی فضایی معرفی کردیم و به طریق مشابه انتخاب یک مدل معتبر برای Σ_t از بخش ساختار وابستگی زمانی، می‌توان مدل‌های معتبر فضایی-زمانی به‌دست آورد (محمدزاده، ۱۳۹۱).

فصل ۲

تقریب لاپلاس آشیانه‌ای جمع‌بسته

مسأله کلیدی در استنباط بیزی، حل انتگرال‌های با بعد بالا است. یک روش قدیمی، تقریب لاپلاس است که تاریخ آن به لاپلاس (۱۷۷۴) برمی‌گردد. این ایده ساده، انتگرال را با بسط تیلور مرتبه دوم تابع زیر انتگرال تقریب می‌زند و آن را به صورت عددی حل می‌کند. با توسعه آشیانه‌ای این روش کلاسیک و ترکیب آن با روش‌های عددی مدرن برای ماتریس‌های تنک^۱، شیوه تقریب لاپلاس آشیانه‌ای جمع‌بسته^۲ توسط رو و همکاران (۲۰۰۹) معرفی شد تا استنباط بیزی را برای مدل‌های گوسی پنهان^۳ تقریب بزند. امروزه LGMها مدل‌های بسیار مفیدی برای تحلیل بیزی به شمار می‌روند و اساس بسیاری از مدل‌های آماری را تشکیل می‌دهند. در این رساله، سعی کرده‌ایم تا از قدرت و شگفتی‌های روش INLA بهره گرفته تا به برازش هرچه کاراتر و در عین حال سریع و ساده مدل‌های چندجمله‌ای بپردازیم. در این فصل به معرفی و نحوه به کارگیری این روش کارا پرداخته‌ایم.

۱.۲ مقدمه‌ای بر INLA

یک مانع کلیدی همیشگی بر سر راه آماردانان و پژوهشگرانی که آمار بیزی را برای تحلیل و مدل‌بندی داده‌های خود انتخاب می‌کنند، اجرای استنباط بیزی است. از دیدگاه ریاضی،

^۱Sparse

^۲Integrated nested Laplace approximation (INLA)

^۳Latent Gaussian models (LGMs)

استنباط بیزی با اصول اولیه زیر تعریف می‌شود:

۱. به کمک اطلاعات موجود در داده‌های مشاهده‌شده، اطلاعات پیشین خود را نسبت به پارامترهای نامعلوم به روز می‌کنیم.
۲. توزیع پسین پارامترهای مدل را به دست می‌آوریم.
۳. آماره‌های مورد نیاز را بر اساس توزیع پسین مانند چندک‌ها، توزیع‌های پسین کناری، میانگین و واریانس محاسبه می‌کنیم.

در عمل، گفتن این اصول آسانتر از انجام آن است. با پیداش و توسعه عملی روش‌های نمونه‌گیری مونت کارلوی زنجیر مارکوفی^۴ (رابرت و کسلا، ۱۹۹۹) پیشرفت‌های چشمگیری در به کارگیری رویکرد بیزی در مسایل مختلف کاربردی ایجاد شد. این روش‌ها یک دستورالعمل کلی برای تولید نمونه‌هایی از توزیع پسین با ساختن زنجیر مارکوف و در نظر گرفتن پسین هدف به عنوان توزیع مانا عرضه می‌کنند. همچنین در سال‌های اخیر با پیدایش و توسعه ابزارهای نرم‌افزاری همچون WinBUGS (اسپیگل‌هاتر، ۱۹۹۵)، JAGS (پلامر، ۲۰۱۶) و Stan (تیم استان، ۲۰۱۵) بر توفیقات این روش‌ها برای استنباط بیزی افزوده شده است. از زمانی که استنباط بیزی به این روش‌ها مجهز شد، آمار بیزی رشد سریعی را در محبوبیت تجربه کرده است به طوری که آمار بیزی هم اکنون در بسیاری از رشته‌ها، شاخه‌های آماری و مجله‌های مربوط به آن‌ها به گستردگی خودنمایی می‌کند. با این وجود و به عقیده بسیاری از کاربران آمار بیزی این انقلاب در آمار بیزی، هنوز نیازمند توجه بیشتری است. چرا که به نظر می‌رسد همچنان با استنباط دست و پاگیری روبرو هستیم که به رایانه‌های پرسرعت و زمانی از کاربر که در رسیدن به مدل مناسب و کارا هدر می‌رود، نیازمند است. حتی اجرای مجدد بسیاری از مدل‌ها را دست و پاگیر و تکرار فرآیند انتخاب مدل را ناممکن کرده است (باکس و تیائو، ۱۹۷۳). بنابراین گاهی استنباط بیزی، به دلایل عملی نبودن به کارگیری آن، مورد استقبال قرار نمی‌گیرد و کنار گذاشته می‌شود. این در حالی است که روش تقریبی INLA، استنباط را در یک فاصله زمانی معقول انجام می‌دهد و در بسیاری از موارد از روش معمول MCMC سریعتر و کاراتر است. اگرچه INLA در عمل به مدل‌های LGM محدود است، اما این رده از مدل‌ها گستره بزرگی از مدل‌های کارآمد و پرمصرف را شامل می‌شوند و افزون بر آن، امروزه بسیاری از مدل‌های آماری را می‌توان براساس LGM‌ها تعریف کرد.

قصه از آن‌جا آغاز می‌شود که رو و همکاران (۲۰۰۹) توانستند نشان دهند با استفاده از تقریب لاپلاس آشیانه‌ای تودرتو می‌توان چگالی‌های کناری پسین را با دقت خیلی خوبی محاسبه کرد. این روش روی مدل‌هایی تمرکز می‌کند که بردار متغیر پنهان در آن‌ها، یک میدان تصادفی گوسی مارکوفی^۵ است. ایده اصلی بر این اساس پایه‌گذاری شد که می‌توان انتگرال توزیع پسین را به صورت حاصلضرب آشیانه‌ای انتگرال‌های با بعد پایین نوشت و سپس

^۴Markov chain Monte Carlo (MCMC)

^۵Gaussian Markov random field (GMRF)

به صورت تقریبی اما با دقت بالا انتگرال‌ها را به صورت عددی حل کرد. روش INLA، در اغلب موارد، با دقتی عمل می‌کند که حتی اختلاف‌های بین نتایج حاصل از آن و نمونه‌گیری MCMC، جزئی و غیرقابل تشخیص هستند. فایده اصلی این تقریب از نظر محاسباتی این است که برآوردهای دقیقی را برای تحلیل‌های بیزی در چند ثانیه یا دقیقه انجام می‌دهد. برای مثال، در مدل‌های معمولی فضایی که بعد n چند هزار است، تقریباً این روش برای همه پسین‌های کناری در یک دقیقه یا چند دقیقه محاسبه می‌شود. در مقابل، الگوریتم‌های MCMC مربوط به آن‌ها ساعت‌ها یا روزها طول می‌کشد تا چگالی‌های کناری پسین را محاسبه کند. همچنین اریبی تقریب خیلی کمتر از خطای MCMC و در عمل ناچیز است. از جمله مهم‌ترین دلایلی که می‌توان برای گسترش فراوان استفاده از روش INLA بیان کرد، فراهم کردن تحلیل‌های بیزی کامل با اجرای سریع و کدنویسی راحت و سریع است بدون آن که نیاز به نوشتن الگوریتم‌های نمونه‌گیری باشد. اما شاید بتوان علت اصلی این گسترش را قابل اجرا بودن برای رده وسیعی از مدل‌ها (مدل‌های گوسی پنهان) دانست چرا که امروزه بسیاری از مدل‌های پویا و مدل‌های فضایی در این رده قرار می‌گیرند.

توسعه نرم‌افزاری این روش، نتیجه چندین سال تلاش و کار سخت بوده که توسط هاوارد رو و تیمش انجام شده است. طی سال‌های ۲۰۰۲ تا ۲۰۰۴، هاوارد به همراه هلد، شروع به درک اهمیت مدل‌هایی را کردند که می‌توان با INLA انجام داد. در سال ۲۰۰۵ آن‌ها کتابی را با عنوان

Gaussian Markov Random Fields: Theory and Applications

در اهمیت این گونه مدل‌ها نوشتند. در سال ۲۰۰۷، اولین اجرای روش INLA در محیط C انجام شد، اما نیاز به فایل‌های دستی ورودی داشت. در سال ۲۰۰۸ دانشجوی دکتری هاوارد به نام سارا مارتینو اولین نسخه آن را برای اجرا در نرم‌افزار R تهیه کرد. برای داشتن مروری بر مثال‌های به کارگرفته شده با INLA می‌توانید به سایت www.r-inla.org و مقاله مروری رو و همکاران (۲۰۱۷) مراجعه کنید. برای درک این روش ابتدا به تعریف مفاهیم زیر می‌پردازیم:

۱. تقریب لاپلاس

۲. مدل‌های گوسی پنهان

۳. میدان‌های تصادفی گوسی مارکوفی

۲.۲ تقریب لاپلاس

تقریب لاپلاس یک روش قدیمی در تقریب انتگرال‌ها به شمار می‌رود. فرض کنید هدف تقریب انتگرال

$$I_n = \int_{\chi} \exp(n f(x)) dx$$

باشد در حالتی که $n \rightarrow \infty$. فرض کنید x_0 نقطه‌ای باشد که $f(x)$ در آن بیشترین مقدار را دارد. بنابراین

$$\begin{aligned} I_n &\approx \int_{\mathcal{X}} \exp \left(n(f(x_0) + \frac{1}{2}(x - x_0)^2 f''(x_0)) \right) dx \\ &= \exp(nf(x_0)) \sqrt{\frac{2\pi}{-nf''(x_0)}} = \tilde{I}_n. \end{aligned}$$

در این جا روش لاپلاس، ایده ساده اما قدرتمند تقریب تابع زیر انتگرال با چگالی توزیع گوسی را انجام می‌دهد. اگر $nf(x)$ را لگاریتم درست‌نمایی و x را پارامتر نامعلوم در نظر بگیریم و قضیه حد مرکزی وقتی $n \rightarrow \infty$ برقرار باشد، آن گاه تقریب گوسی دقیق است. برای انتگرال‌های با بعد بالا، خطای تقریب $O(n^{-1})$ است. بنابراین

$$I_n = \tilde{I}_n(1 + O(n^{-1})).$$

این یک نتیجه خوب است، چرا که خطای نسبی از مرتبه n^{-1} در مقابل $n^{-\frac{1}{2}}$ در به کارگرفتن روش‌های معمول استنباط بر پایه شبیه‌سازی است.

تقریب لاپلاس یک ابزار کلیدی در زمان‌های قبل از عمومی شدن الگوریتم‌های MCMC به شمار می‌رود. فرض کنید بخواهیم توزیع کناری $\pi(\gamma_1)$ را از توزیع توأم $\pi(\gamma)$ به دست آوریم، که در آن $\gamma = (\gamma_1, \dots, \gamma_n)'$ می‌توان نوشت:

$$\begin{aligned} \pi(\gamma_1) &= \frac{\pi(\gamma)}{\pi(\gamma_{-1}|\gamma_1)} \\ &\approx \frac{\pi(\gamma)}{\pi_G(\gamma_{-1}; \mu(\gamma_1), Q(\gamma_1))} \Big|_{\gamma_{-1}=\mu(\gamma_1)} \end{aligned} \quad (1.2)$$

در واقع $\pi(\gamma_{-1}|\gamma_1)$ با چگالی توزیع گوسی تقریب زده می‌شود که در آن $\pi_G(\cdot; \mu, Q)$ تقریب گوسی $\pi(\cdot)$ است. تیرنی و کادان (۱۹۸۶) نشان دادند اگر $\pi(\gamma) \propto \exp(nf_n(\gamma))$ یعنی اگر $f_n(\gamma)$ میانگین لگاریتم درست‌نمایی باشد، خطای نسبی تقریب (۱.۲) برابر است با $O(n^{-\frac{1}{2}})$.

۱.۲.۲ مدل‌های گوسی پنهان

مدل‌های گوسی پنهان را به کمک چارچوب مدل‌های سلسله‌مراتبی توضیح می‌دهیم. با فرض استقلال مشاهدات $y = (y_1, \dots, y_n)'$ با تابع (چگالی) احتمال $p(\cdot|\theta_1)$ ، به شرط میدان تصادفی گوسی تصادفی x و بردار ابرپارامتر θ_1 ، می‌توان نوشت

$$p(y|x, \theta_1) \sim \prod_{i=1}^n p(y_i|x_i, \theta_1)$$

به‌طوری که

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

بردار x شامل همه قسمت‌های تصادفی مدل که نشان‌دهنده ساختار وابستگی موجود در مدل هستند، می‌شود. بنابراین درست‌نمایی دارای ابرپارامترهای $\theta' = (\theta'_1, \theta'_2)$ است. بنابراین، توزیع پسین مدل به صورت

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

نتیجه می‌شود. توجه داشته باشید که x توزیع توام مربوط به همه پارامترهای موجود در پیشگوی خطی است و θ بردار ابرپارامترهای مدل که لزوماً دارای توزیع گوسی نیست. بنابراین، این مدل یک LGM است اگر و تنها اگر توزیع توام پارامترهای مدل گوسی باشد نه ابرپارامترها. در واقع مدل‌های گوسی پنهان همان مدل‌های STAR هستند که اثرات ساختاری آن‌ها گوسی است. در حالت کلی برای کارایی روش تقریب INLA در مدل‌های گوسی پنهان، باید شرایط زیر برقرار باشند:

۱. تعداد ابرپارامترها کم و معمولاً بین ۲ تا ۵ است و نباید از ۲۰ بیشتر شود.
 ۲. توزیع میدان پنهان $x | \theta$ گوسی و با ویژگی مارکوفی است و n می‌تواند بزرگ بین 10^3 تا 10^5 باشد.
 ۳. داده‌های y ، به شرط داشتن x و θ ، مستقل شرطی هستند. در نتیجه مشاهدات y_i به عناصر میدان پنهان یعنی x_i ها بستگی دارند.
- این شرایط هم به دلایل محاسباتی و هم به خاطر ایجاد اطمینان بالا از داشتن یک تقریب دقیق، ضروری است.

۲.۲.۲ میدان‌های تصادفی گوسی مارکوفی

همان‌طور که پیشتر بیان شد، یکی از شرایط لازم برای کارایی روش INLA مارکوفی بودن میدان تصادفی x است. فرض کنید بردار تصادفی $x = (x_1, x_2, \dots, x_n)'$ دارای توزیع نرمال چندمتغیره با میانگین μ و ماتریس دقت معین مثبت Q با تابع چگالی

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

است. بردار x دارای ویژگی مارکوفی است، هرگاه x_i و x_j به شرط سایر اعضای باقیمانده x_{-ij} ، برای تعداد کمی از $\{i, j\}$ ها از یکدیگر مستقل باشند. اهمیت این ویژگی در آن است که اگر دو عضو x_i و x_j مستقل شرطی باشند، آن‌گاه درایه متناظر با آن‌ها در ماتریس دقت صفر خواهد بود. یعنی $Q_{ij} = 0$. نتیجه عملی مارکوفی بودن میدان تصادفی آن است که ماتریس دقت Q

تنک خواهد شد و روش‌های بسیار کارایی برای تجزیه ماتریس‌های تنک وجود دارند (رو و هلد، ۲۰۰۵) که باعث افزایش سرعت محاسبات می‌شود.

بنابراین در عمل برای داشتن کارایی محاسباتی، میدان تصادفی پنهان نه تنها گوسی بلکه باید مارکوفی (تنک) هم باشد (رو و هلد، ۲۰۰۵؛ هلد و رو ۲۰۱۰). در واقع یک GMRF یک میدان گوسی با ویژگی اضافه‌تر استقلال شرطی است. شاید ساده‌ترین مثال از یک مدل مارکوفی، مدل اتورگرسیو مرتبه اول با معادله

$$x_t = \phi x_{t-1} + \epsilon_t, \quad t = 1, 2, \dots, m,$$

باشد. در این مدل، همبستگی بین x_s و x_t برابر است با $\phi^{|s-t|}$ و ماتریس کوواریانس $m \times m$ چگال است. با این وجود، x_s و x_t برای $|s-t| > 1$ به شرط سایر x ها از هم مستقل‌اند و بنابراین ماتریس دقت حاصل، یک ماتریس تنک سه قطری است. به همین دلیل است که در ادبیات INLA در میدان گوسی از ماتریس دقت برای نمایش توزیع استفاده می‌شود تا ماتریس کوواریانس.

در نظر گرفتن GMRF در محاسبات خیلی کاراتر و بهتر است زیرا در این صورت به جای یک ماتریس چگال که هزینه محاسباتی بالایی دارد، با یک ماتریس تنک سروکار داریم. در مثال اتورگرسیو، ماتریس دقت، یک ماتریس سه قطری است که تجزیه آن به زمان محاسباتی $O(m)$ نیاز دارد. در حالی که در حالت کلی تجزیه یک ماتریس چگال $m \times m$ به محاسبات با مرتبه $O(m^3)$ نیاز دارد. در مدل‌های با ساختار فضایی نیز زمان محاسبات به $O(m^{3/2})$ و نیاز به حافظه به $O(m \log(m))$ کاهش می‌یابد.

۳.۲ مراحل روش تقریب INLA

فرض کنید بردار $y = (y_1, \dots, y_n)'$ ، n مشاهده از توزیع پاسخ با خانواده توزیع‌های نمایی با میانگین μ_i باشد. همچنین فرض کنید پیشگوی جمعی-ساختاری مدل نسبت به میانگین، تابعی از متغیرهای تبیینی باشد. توجه داشته باشید که پیشگوی ساختاری می‌تواند تابعی از اثرات ساختاری مانند توابع ناخطی، اثر تصادفی و اثر فضایی نیز باشد، که این جملات بردار پنهان x را می‌سازند. توزیع متغیر پاسخ به بردار ابرپارامترهای $\theta = (\theta_1, \theta_2)$ نیز وابسته است. مشاهدات پاسخ به شرط بردار متغیرهای پنهان x از یکدیگر مستقل هستند. بنابراین تابع درستنمایی مدل به صورت

$$p(y|x, \theta) = \prod_{i \in I} p(y_i|x_i, \theta),$$

است، که در آن $I \subseteq \{1, \dots, n\}$ نشان‌دهنده اندیس مشاهدات پاسخ است. در استنباط بیزی، هدف به‌دست آوردن توزیع پسین پارامترهای مدل و ابرپارامترها است که روش به‌دست آوردن آن در INLA در سه گام اجرا می‌شود:

گام اول

برای توزیع شرطی کامل

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

با استفاده از بسط تیلور $\log p(y_i|x_i, \boldsymbol{\theta})$ حول مد توزیع شرطی کامل $\mathbf{x}$ ، یعنی

$$\mathbf{x}^*(\boldsymbol{\theta}) = (x_1^*(\boldsymbol{\theta}), \dots, x_{n_d}^*(\boldsymbol{\theta}))',$$

که در آن n_d بعد بردار پنهان $\mathbf{x}$ است، تقریب گوسی به صورت

$$\begin{aligned} \tilde{\pi}_G(\mathbf{x}|\mathbf{y}, \boldsymbol{\theta}) &\propto \exp\left(-\frac{1}{2}\mathbf{x}'\mathbf{Q}\mathbf{x} + \sum_{i \in I} g_i(x_i)\right) \\ &\propto \exp\left(-\frac{1}{2}\mathbf{x}'\mathbf{Q}\mathbf{x} + \sum_{i \in I} (a_i + b_i x_i - \frac{1}{2} c_i x_i^2)\right), \end{aligned}$$

به دست می آید، که در آن $c_i = -g''(x_i^*(\boldsymbol{\theta}))$ و $b_i = g'_i(x_i^*(\boldsymbol{\theta})) + x_i^*(\boldsymbol{\theta})c_i$ و جمله a_i را می توان از تناسب حذف کرد (رو و همکاران، ۲۰۰۵). سپس تقریب لاپلاس چگالی پسینی

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

با بسط مخرج کسر حول $\mathbf{x} = \mathbf{x}^*(\boldsymbol{\theta})$ به صورت

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

محاسبه می شود.

گام دوم

تقریب لاپلاس چگالی شرطی

$$\pi(x_i|\boldsymbol{\theta}, \mathbf{y}) = \frac{\pi(\mathbf{x}, \boldsymbol{\theta}|\mathbf{y})}{\pi(\mathbf{x}_{-i}|x_i, \boldsymbol{\theta}, \mathbf{y})},$$

با بسط مخرج کسر حول $\mathbf{x}_{-i} = \mathbf{x}_{-i}^*(x_i, \boldsymbol{\theta})$ به صورت

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

محاسبه می شود، که در آن $\mathbf{x}_{-i}^*(x_i, \boldsymbol{\theta})$ مد چگالی $\pi(x_i|x_i, \boldsymbol{\theta}, \mathbf{y})$ و $\tilde{\pi}_{GG}(x_i|x_i, \boldsymbol{\theta}, \mathbf{y})$ تقریب گوسی $\pi(x_i|x_i, \boldsymbol{\theta}, \mathbf{y})$ است. همچنین $\mathbf{x}_{-i}$ بردار حاصل از حذف درایه i ام $\mathbf{x}$ است.

گام سوم

در نهایت جمع‌های متناهی

$$\begin{aligned}\tilde{\pi}(x_i|\mathbf{y}) &= \sum_{k=1}^{n_{p_1}} \tilde{\pi}(x_i|\boldsymbol{\theta}_k, \mathbf{y}) \tilde{\pi}(\boldsymbol{\theta}_k|\mathbf{y}) \Delta_k; \quad i = 1, \dots, n_d, \\ \tilde{\pi}(\theta_j|\mathbf{y}) &= \sum_{q=1}^{n_{p_2}} \tilde{\pi}(\theta_q|\mathbf{y}) \Delta_{jq}; \quad j = 1, \dots, \ell\end{aligned}$$

تقریب زده می‌شوند، که در آن n_{p_1} و n_{p_2} تعداد $\boldsymbol{\theta}$ های انتخاب‌شده از تکیه‌گاه ابرپارامترها است، ℓ بعد بردار ابرپارامترها است و وزن‌های Δ_k و Δ_{jq} با استفاده از رو و همکاران (۲۰۰۹) به صورت

$$\begin{aligned}\Delta_k &= \{(n_{p_1})(f_0^\gamma - 1)[1 + \exp(\frac{-mf_0^\gamma}{\gamma})]\}^{-1} \\ \Delta_{jq} &= \{(n_{p_2})(f_0^\gamma - 1)[1 + \exp(\frac{-(m-1)h_0^\gamma}{\gamma})]\}^{-1},\end{aligned}$$

محاسبه می‌شوند که در آن f_0 و h_0 هر مقدار ثابت بزرگتر از یک و m حاصل انتگرال $\boldsymbol{\theta}'\boldsymbol{\theta}$ است. در واقع این جمع‌ها به ترتیب تقریب چگالی‌های پسین کناری

$$\begin{aligned}\pi(x_i|\mathbf{y}) &= \int \pi(x_i|\mathbf{y}, \boldsymbol{\theta}) \pi(\boldsymbol{\theta}|\mathbf{y}) d\boldsymbol{\theta}, \quad i = 1, \dots, n_d, \\ \pi(\theta_j|\mathbf{y}) &= \int \pi(\boldsymbol{\theta}|\mathbf{y}) d\boldsymbol{\theta}_{-j}, \quad j = 1, \dots, \ell\end{aligned}$$

هستند، که در آن $\boldsymbol{\theta}_{-j}$ بردار حاصل از حذف درایه j ام $\boldsymbol{\theta}$ است.

به‌طور خلاصه، رویکرد INLA بر اساس ساختار رده LGM طراحی شده است؛ بعد $\boldsymbol{\theta}$ پایین است؛ $x|\boldsymbol{\theta}$ یک GMRF است و مشاهدات پاسخ مستقل شرطی هستند، به‌طوری که y_i تنها به x_i و $\boldsymbol{\theta}$ بستگی دارد. برای فهم بیشتر از فرآیند تقریب، مثال زیر از رو و همکاران (۲۰۱۷) را تشریح می‌کنیم.

مثال ۱.۳.۲. مدل

$$\begin{aligned}y_i|\lambda_i &\sim \text{Pois}(\lambda_i), \quad i = 1, \dots, n, \\ \log(\lambda_i) &= \eta_i = g(\beta)u_{j(i)},\end{aligned}$$

را در نظر بگیرید، که در آن $\beta \sim N(0, 1)$ ، $\mathbf{u} \sim N(0, \mathbf{Q}^{-1})$ و $g(\cdot)$ یک تابع یکنوا است. اندیس $j(i)$ طوری ساخته شده است که بعد $\mathbf{u}$ ثابت باشد و به n بستگی نداشته باشد و u_j ها به همان تعداد دفعه مشاهده شده‌اند. محاسبه پسین کناری β و u_j با مشکل روبرو است، زیرا ضرب دو مدل گوسی و ناگوسی را داریم. راهکار INLA این است که این مسئله را به چند مسئله کوچکتر که تقریباً گوسی هستند، تبدیل کند و در نهایت تقریب لاپلاس را برای آن‌ها به کار

گیرد. ایده کلیدی در این جا، نوشتن توزیع شرطی است. برای این مسئله توزیع شرطی β به صورت

$$\pi(\beta|\mathbf{y}) \propto \pi(\beta) \int \prod_{i=1}^n p(y_i|\lambda_i) \times p(\mathbf{u}) d\mathbf{u}, \quad (2.2)$$

است که در آن $\lambda_i = \exp(g(\beta)u_{j(i)})$. برای تقریب خوب انتگرال، باید تابع زیر انتگرال به چگالی گوسی نزدیک باشد. توزیع پسین کناری u_j نیز به صورت

$$\pi(u_j|\mathbf{y}) = \int \pi(u_j|\beta, \mathbf{y}) \times \pi(\beta|\mathbf{y}) d\beta. \quad (3.2)$$

است. توجه داشته باشید که در این جا می توان انتگرال را به طور مستقیم حساب کرد، زیرا β یک بعدی است. مشابه با (2.2) می توان نوشت

$$\pi(\mathbf{u}|\beta, \mathbf{y}) \propto \prod_{i=1}^n p(y_i|\lambda_i) \times p(\mathbf{u}) \quad (4.2)$$

که باید نزدیک به چگالی گوسی باشد. تقریب زدن $\pi(u_j|\beta, \mathbf{y})$ ، همان تقریب انتگرال مربوط به $\pi(\mathbf{u}|\beta, \mathbf{y})$ اما با یک بعد کمتر است، زیرا u_j ثابت است که باز هم به چگالی گوسی نزدیک است.

یک نکته کلیدی این است که می توان مسئله را به سه مسئله ریزتر زیر تبدیل کرد:

۱. تقریب $\pi(\beta|\mathbf{y})$ به کمک (2.2).

۲. تقریب $\pi(u_j|\beta, \mathbf{y})$ برای همه j ها و β های دلخواه.

۳. محاسبه $\pi(u_j|\mathbf{y})$ با استفاده از ترکیب دو تقریب در مراحل ۱ و ۲ در ترکیب با (4.2).

اگرچه در ظاهر هزینه ای که پرداخت می کنیم افزایش می یابد و مسئله پیچیده تر می شود، زیرا باید مرحله دوم را برای تمام مقادیر β ها حساب کنیم، اما با کمی تأمل درخواهیم یافت که با جایگزین کردن وابستگی های پیچیده با شرطی کردن و انتگرال های عددی، مسئله به تقریب لاپلاس توزیع هایی که به گوسی نزدیک هستند منتهی می شود.

حال یک سؤال مهم این است که آیا می توان این اصول را برای LGMها به کار گرفت به طوری که از دقت و کارایی لازم برخوردار باشد؟ پاسخ، مثبت است. می توان اصول بالا را برای مدل های LGM نیز به کار گرفت به طوری که β را با θ و u را با x جایگزین کرد. تقریبی که در نتیجه آن به دست می آید، هر چند از محاسبات سریعی برخوردار است اما کمی دقت خود را از دست می دهد.

۴.۲ مثال‌های بیشتر برای مقایسه

برای مقایسه سرعت تقریب حاصل از روش INLA با سایر نرم‌افزارهای استنباط بیزی، دو مثال را تشریح می‌کنیم.

مثال ۱.۴.۲. در این مثال، مشاهدات پاسخ را از مدل رگرسیون خطی ساده

$$y_i = \alpha + \beta x_i + \epsilon_i, \quad i = 1, \dots, n,$$

شبیه‌سازی کردیم، که در آن ϵ_i دارای توزیع نرمال استاندارد با دقت σ^{-1} است. متغیر تبیینی x را از یک توزیع نرمال با میانگین ۶ و واریانس ۲ تولید کردیم و مقادیر واقعی پارامترهای رگرسیونی با تولید از یک توزیع نرمال استاندارد، $\beta = -0.674$ و $\alpha = 1.09$ انتخاب شدند. مقدار واقعی σ^{-1} را نیز برابر با یک در نظر گرفتیم. برای مقایسه نتایج حاصل از رهیافت INLA، مدل رگرسیون بیزی را با همان پیشین‌های نرمال در روش INLA برای ضرایب رگرسیونی α و β و پیشین گاما برای پارامتر دقت، با ابرپارامترهای یکسان، به زبان برنامه‌نویسی JAGS در این نرم‌افزار معرفی و شبیه‌سازی از آن را اجرا کردیم. برآوردهای میانگین پسینی پارامترها توسط هر دو روش INLA و JAGS برای $n = 100$ در جدول ۱.۲ گزارش شده‌اند. برآوردها با هر دو روش به مقادیر واقعی نزدیک هستند. مدت زمان اجرای مدل با دو روش برای اندازه‌های نمونه برابر با ۱۰۰، ۵۰۰، ۵۰۰۰، ۲۵۰۰۰ در جدول ۲.۲ گزارش شده‌اند. از روی نتایج این جدول، به روشنی سرعت بالا و کارایی محاسباتی روش INLA مشهود است.

جدول ۱.۲: نتایج برازش مدل رگرسیون ساده

INLA	JAGS	مقدار واقعی	
-۰/۶۷۱	-۰/۶۶۹	-۰/۶۷۴	α
۱/۰۶	۱/۰۰	۱/۰۹	β

جدول ۲.۲: مقایسه زمان اجرا برای مدل رگرسیون ساده

INLA	JAGS	n
۰/۱۷۶	۴/۱۹	۱۰۰
۱/۲۴۳	۱۸/۱۴۱	۵۰۰
۲/۷۸۷	۳۸۱/۵۷۳	۵۰۰۰
۱۳/۲۷	۲۲۰۳/۶۷۹	۲۵۰۰۰
۵۲/۷۸۷	۸۸۷۳/۸۳۶	۱۰۰۰۰۰

مثال ۲.۴.۲. معمولاً برای نمایش همزمان وابستگی اندازه دور سر جنین به سن مادر و تفاوت بین رشد دور سر در جنین‌ها، از مدل عرض از مبدا و شیب تصادفی

$$y_{ft} = (\beta_0 + u_{0f}) + (\beta_1 + u_{1f})\text{tga} + e_{ft}, \quad u_{0f} \sim N(0, \epsilon), \quad u_{1f} \sim N(0, \epsilon_1), \quad e_{ft} \sim N(0, \sigma),$$

استفاده می‌کنند، که در آن y_{ft} پاسخ طولی (دور سر جنین) برای جنین f در زمان t و tga سن مقیاس‌بندی‌شده مادر است.

جدول ۳.۲: نمایی از ساختار داده‌های مربوط به مثال اندازه دور سر جنین

شماره مشاهده	id	y	ga	tga
۱	۱	۲۱۱	۲۳/۰۰	۱۶/۸۳
۲	۱	۲۷۴	۲۸/۴۳	۱۹/۰۰
۳	۱	۳۱۴	۳۳/۴۳	۲۰/۳۹
.	.	.	.	.
.	.	.	.	.
.	.	.	.	.
۳۰۹۵	۷۰۷	۲۹۲	۲۹/۵۴	۱۹/۳۰
۳۰۹۶	۷۰۷	۳۴۸	۳۷/۵۷	۲۱/۱۱

نمایی از ساختار داده‌ها در جدول ۴.۲ قابل مشاهده است، که در آن id کد شناسایی مادران، y اندازه دور سر جنین (پاسخ)، ga سن مادر و tga همان سن مادر است که برای راحتی برازش در ۱۹/۵۰ سال مرکزی شده است. برای دسترسی کامل به این مجموعه داده می‌توانید به لینک <http://BendixCarstensen.com/Bayes/Cph-2012/data/fetal.csv> مراجعه کنید. به منظور مقایسه روش INLA با یک روش مبتنی بر درست‌نمایی کلاسیک و MCMC، به کمک بسته‌های نرم‌افزاری lme4 و JAGS مدل را به داده‌ها برازش دادیم و نتایج را با نتایج حاصل از روش INLA مقایسه کردیم. برآوردهای متناظر برای اثر ثابت، پارامتر دقت اثرات تصادفی، انحراف استاندارد جمله باقیمانده و همبستگی بین شیب و عرض از مبدا، به ترتیب در جدول‌های ۴.۲، ۵.۲، ۶.۲ و ۷.۲ گزارش شده‌اند.

جدول ۴.۲: برآوردهای عرض از مبدا و شیب در مدل

	INLA	JAGS	lme4
عرض از مبدا	۲۸۶/۲۹	۲۸۵/۵۸	۲۸۶/۲۷
tga	۲۶/۴۸	۲۶/۳۴	۲۶/۴۸

جدول ۵.۲: برآوردهای پارامترهای دقت اثرات تصادفی

	INLA	JAGS	lme4
عرض از مبدا	۶۷/۹۳	۷۶/۲۳	۷۴/۴۹
tga	۱/۰۰۲	۱/۱۵	۱/۲۱

جدول ۶.۲: برآورد پارامتر دقت مشاهدات

INLA	JAGS	lme4
۸/۶۴	۸/۶۳	۸/۵۸

جدول ۷.۲: برآورد همبستگی بین شیب و عرض از مبدأ

INLA	JAGS	lme4
۰/۷۸	۰/۷۶	۰/۷۴

در تمام نتایج، هر سه روش عملکرد مشابهی دارند و نمی‌توان تفاوت چشمگیری بین آن‌ها تعیین کرد. مدت زمان محاسبات برای JAGS و INLA به ترتیب ۱۱۴/۵۹ و ۱۴/۳ دقیقه است، که به روشنی کارایی روش INLA قابل درک است.

با همه ویژگی‌های خوبی که روش INLA دارد، متأسفانه به دلیل عدم امکان تعبیه هسته توزیع چندجمله‌ای در بسته R-INLA نمی‌توان پاسخ‌های رسته‌ای اسمی را با آن مدل‌بندی و تحلیل کرد. بنابراین، به کمک تبدیلی به نام تبدیل چندجمله‌ای-پوآسن^۶ (بیکر، ۱۹۹۴)، ابتدا هسته چندجمله‌ای را به پوآسن تبدیل کرده و سپس به کمک روش INLA به تحلیل پاسخ‌های رسته‌ای می‌پردازیم. در فصل چهارم این تبدیل را تشریح می‌کنیم.

در این فصل سعی کردیم تا نشان دهیم روش INLA یک تکنیک برای اجرایی کردن برآزش مدل‌های پیچیده است و نه حتی یک روش شبیه‌سازی مانند MCMC و نه صرفاً یک روش کدنویسی، بلکه روشی است که کارایی محاسباتی به همراه سادگی نسبی کدنویسی در آن موجب شده است تاکنون افراد زیادی برای به کارگیری مدل‌های فضایی یا فضایی-زمانی پیچیده یا هر مدل پیچیده دیگری که در چارچوب رده LGM قرار گیرد، با کاربردهای مختلف و فراوان در آن اقدام کنند.

^۶ Multinomial-Poisson transformation (MPT)

فصل ۳

شناسایی پذیری

شناسایی پذیری^۱ یکی از پیش نیازهای اساسی در مدل‌های آماری به شمار می‌رود و استنباط بدون شناسایی پذیری بی‌معنی است. بنابراین، لازم است قبل از برازش مدل شناسا بودن آن را مورد بررسی قرار دهیم. در این فصل مفهوم شناسایی پذیری در مدل‌ها را بررسی و تشریح می‌کنیم. پس از درک مفاهیم و شناخت راه‌حل‌های پیش رو برای رفع مشکل عدم شناسایی پذیری، روش پیشنهادی برای رفع آن در مدل‌های چندجمله‌ای با پیشگوی جمعی-ساختاری را در مدل مشابه آن به چالش می‌کشیم. در انتها برخی نتایج نظری را بیان می‌کنیم.

۱.۳ مقدمه

استنباط آماری به ما می‌آموزد تا ”چگونه” از داده‌ها بیاموزیم در حالی که تحلیل شناسایی ”چیزی” را که ما می‌توانیم از داده‌ها بیاموزیم، توضیح می‌دهد. اگرچه از نظر منطقی ”چه چیزی” بر ”چگونه” مقدم است، اما متأسفانه مفهوم شناسایی پذیری توجه نسبتاً کمی را در جامعه آماری به خود اختصاص داده است و ما باید نسبت به مفاهیمی از آمار که به امکان پذیر بودن تشخیص منحصر به فرد پارامترها توسط توزیع یا متغیر تصادفی مشاهده شده می‌پردازند، توجه بیشتری داشته باشیم (باسه و بوجینو، ۲۰۲۰). مسئله شناسایی پذیری در بسیاری از علوم مانند اقتصادسنجی و زیست‌شناسی که افراد در آن مجبور به مدل‌بندی ساختارهای

^۱Identifiability

موجود در طبیعت هستند، مطرح است. این گوناگونی و گستردگی موجب شده است تا بسیاری از افرادی که رشته اصلی آن‌ها آمار نبوده، در پی نیاز به مدل‌بندی، به توضیح و مفهوم‌سازی مناسب برای این مسئله روی آورند. تاریخ اولین پژوهش‌ها در این زمینه را به دانشمندان علم اقتصاد به حداقل قبل از سال ۱۹۳۰ نسبت می‌دهند. اولین بار کوپمن (۱۹۴۹) مفهوم شناسایی‌پذیری را مطرح کرد. بعد از آن هرویچ (۱۹۵۰) و کوپمن و ریرسول (۱۹۵۰) به تعریف ریاضی این مفهوم پرداخته و نظریه مربوط به آن را در مدل‌های آماری توسعه دادند. به دنبال آن افرادی از رشته‌های مختلف برای مدل‌بندی سیستم‌های بیولوژی، بوم‌شناسی، مدل‌های پارامتری و ناپارامتری مطرح شدند. این افراد به ترتیب شامل گودمن (۱۹۵۹)، رتنبرگ (۱۹۷۱)، هیسایو (۱۹۸۳)، جاکز و پری (۱۹۹۰)، پولانیو و دپراگانکا بریرا (۱۹۹۴)، کروس و مانسکی (۲۰۰۲)، ماتزکین (۲۰۰۷) و مانسکی (۲۰۰۹) است. این واگرایی و فقدان نظریه یکپارچه، موجب شده است تا این بخش از علم آمار توسعه جدی را به خود نبیند.

در منظر علم آمار، مدل مفید بر سه پایه استوار است (وایلند و همکاران، ۲۰۲۱). یک مدل خوب در درجه اول باید با یک دقت منطقی داده‌ها را توصیف کند. دوم این که باید قدرت پیش‌گویی بالایی داشته باشد. مدل‌هایی که دارای این دو ویژگی باشند، خوب محسوب می‌شوند. در درجه سوم یک مدل باید توانایی برازش شدن به سیستم‌های بیولوژیکی را داشته باشد.

در یک فرآیند مدل‌بندی، معمولاً افراد از ساده‌ترین مدل شروع می‌کنند که توانایی توصیف داده‌ها را ندارد و بد تلقی می‌شود. افزودن همزمان پیچیدگی به مدل به گونه‌ای که قابل برازش باشد و توانایی توصیف داده را داشته باشد، نیازمند توجه به مسئله قابل شناسایی بودن مدل است. یعنی مدل هم از قدرت پیش‌گویی بالایی برخوردار بوده و هم توانایی تعیین پارامترها به کمک داده‌ها را داشته باشد که همان شناسایی‌پذیری مدل است. مفهوم شناسایی‌پذیری با قدرت زیاد به انتقال مدل از بد به خوب مربوط می‌شود. به عبارتی دیگر، تحلیل شناسایی برای داشتن یک مدل خوب، ضروری است.

یکی از اهداف ما در این رساله افزایش سطح شناساگر بودن مدل لوجیت چندجمله‌ای است، به گونه‌ای که همزمان از ساختارهای پیچیده فضایی و فضایی-زمانی در یک چارچوب بیزی به کمک روش INLA بهره گرفته باشد.

۲.۳ شناسایی‌پذیری مدل

فرض کنید مدل آماری (پارامتری) $(\mathcal{Y}, \varphi, P)$ را داشته باشیم که در آن $\mathcal{Y}$ فضای نمونه، φ سیگما-میدان تعریف‌شده روی آن و P خانواده‌ای از اندازه احتمال روی $(\mathcal{Y}, \varphi)$ است که با $P = \{P_\theta : \theta \in \Theta\}$ مشخص شده و Θ فضای پارامتر است.

تعریف ۱.۲.۳. (پائولینو و پریرا، ۱۹۹۴). دو نقطه از Θ از نظر مشاهده معادل هستند، اگر به

ازای هر $A \in \wp$ ، $P_{\vartheta_0}(A) = P_{\vartheta_1}(A)$ ، در این صورت می‌توان نوشت $\vartheta_0 \sim \vartheta_1$.

تعریف ۲.۲.۳. (اسوارتز و همکاران، ۲۰۰۴). فرض کنید $F = \{F_\theta, \theta \in \Theta\}$ خانواده‌ای از توابع توزیع باشد. همچنین فرض کنید A متغیر تصادفی قابل مشاهده^۲ از این خانواده باشد. اگر حداقل یک جفت $(\theta, \theta'), \theta \neq \theta', \theta, \theta' \in \Theta$ وجود داشته باشد که به ازای هر a از تکیه‌گاه متغیر $F_\theta(a) = F_{\theta'}(a)$ ، می‌توان گفت θ با A قابل شناسایی نیست.

در واقع، بنا به تعریف ۲.۲.۳، در این حالت تابع درست‌نمایی برای همه داده‌ها ثابت است.

تعریف ۳.۲.۳. (اسوارتز و همکاران، ۲۰۰۴). اگرچه مفهوم نزدیک ناشناس بودن^۳ از قدرت کمتری برخوردار است، ولی در اشاره به تعریف قبل می‌توان گفت θ تقریباً ناشناس است اگر برای همه مقادیر m در تکیه‌گاه متغیر M ، $F_\theta(m) \approx F_{\theta'}(m)$.

در یک کلام می‌توان گفت ناشناس بودن مدل به داده‌ها اجازه نمی‌دهد تا تفاوتی بین مقادیر پارامترها تشخیص داده شود.

مثال ۱.۲.۳. فرض کنید X دارای توزیع نرمال با میانگین $\mu_1 - \mu_2$ باشد. بنابراین $\mu_1 - \mu_2$ را می‌توان با مشاهداتی از X برآورد کرد، اما پارامترهای μ_1 و μ_2 به‌طور منحصر به فرد قابل شناسایی و برآورد نیستند. در حقیقت می‌توان به تعداد نامتناهی از جفت‌های

$$\{(\mu_i, \mu_j), i, j = 1, 2, \dots (i \neq j),\}$$

فکر کرد به‌گونه‌ای که $\mu_i - \mu_j = \mu_1 - \mu_2$. فقط در صورتی می‌توان μ_1 و μ_2 را به‌طور منحصر به فرد برآورد کرد که قابل شناسایی باشند.

همان‌طور که در مثال بالا قابل مشاهده است، در بسیاری از مسائل آماری فرض می‌شود توزیع آماری مورد مطالعه به‌جز برای مجموعه‌ای از پارامترها کاملاً معلوم است. بدیهی است قبل از هر گونه استنباط باید مدل از نظر شناسایی پذیری مورد بررسی قرار گیرد. در این صورت مطمئن می‌شویم که متغیر تصادفی مورد نظر از توزیع‌های متفاوت با پارامترهای متفاوت پیروی نمی‌کند. گاهی اوقات ممکن است پارامتر θ خود قابل شناسایی نباشد، اما یک تابع ثابت از آن مانند $\varphi(\theta)$ قابل شناسایی باشد. یعنی برای هر θ و θ' که متعلق به Θ باشد، به ازای تمام مقادیر a از تکیه‌گاه A ، $F_\theta(a) = F_{\theta'}(a)$ ما را به $\varphi(\theta) = \varphi(\theta')$ می‌رساند. در این صورت می‌توان گفت θ به‌طور جزئی شناسایی پذیر^۴ است. در مثال قبل F_θ یک خانواده از توزیع‌های نرمال، $\theta = (\mu_1, \mu_2)$ و $\varphi(\theta) = \mu_1 - \mu_2$ است.

برای تبدیل θ به یک متغیر شناسایی پذیر، می‌توان متغیر تصادفی Z را به مدل اضافه کرد به‌طوری که θ توسط متغیر تصادفی تقویت شده (A, Z) قابل شناسایی باشد. در این میان خود

^۲ Observable

^۳ Near nonidentifiability

^۴ Partially identifiable

متغیر تصادفی Z ممکن است نتواند θ را تشخیص دهد. در این صورت مسئله ناشناس بودن قابل اصلاح^۵ است. در مثال زیر این موضوع را تشریح می‌کنیم.

مثال ۲.۲.۳. فرض کنید X_1 و X_2 دو متغیر تصادفی مستقل از توزیع نمایی با تابع توزیع

$$F_i(x) = 1 - \exp(-\lambda_i x), \quad \lambda_i > 0, \quad x > 0 \quad i = 1, 2,$$

باشند. به سادگی می‌توان دید $U = \min(X_1, X_2)$ دارای توزیع نمایی با پارامتر $\lambda_1 + \lambda_2$ است. در این جا λ_1 و λ_2 با U قابل شناسایی نیستند، زیرا تعداد نامتناهی از جفت متغیرهای تصادفی مستقل با توزیع نمایی وجود دارند که مجموع پارامترهای آن‌ها می‌تواند $\lambda_1 + \lambda_2$ باشد. در مقابل، $\theta = (\lambda_1, \lambda_2)$ به‌طور جزئی قابل شناسایی است، زیرا $\varphi(\theta) = \lambda_1 + \lambda_2$ همیشه قابل شناسایی است.

برمن (۱۹۶۳) نشان داد اگر X_i ها به‌طور مستقل توزیع شده باشند و بتوان می‌نیم را شناسایی کرد، آن‌گاه می‌توان توابع توزیع F_i را به‌طور منحصر به فرد بر اساس توابع یکنوا

$$H_k(x) = P(U \leq x, I = k), \quad k = 1, 2, \dots, l,$$

که در آن I پارامتر اضافه است و شناسایی‌پذیری مدل را اصلاح می‌کند، تعیین کرد. به‌عنوان مثال، می‌توان قالب

$$F_k(x) = 1 - \exp\left\{-\int_{-\infty}^x \left(1 - \sum_{j=1}^l H_j(t)\right)^{-1} dH_k(t)\right\}, \quad k = 1, 2, \dots, l,$$

را در نظر گرفت. نتایج برمن در مثال قبل نشان می‌دهد که مسئله ناشناس بودن وقتی از (U, I) استفاده کنیم، قابل اصلاح است.

اگر پارامترها برآورد نشوند یا یک راه منحصر به فرد برای برآورد آن‌ها وجود نداشته باشد، آن‌گاه گوییم مسئله ناشناس^۶ است. اگر یک راه منحصر به فرد برای شناسایی پارامترها وجود داشته باشد، آن‌گاه گوییم پارامترها شناسایی^۷ می‌شوند. اگر پارامترها در بیشتر از یک راه مستقل شناسایی شوند، گوییم پارامترها بیش‌شناس^۸ هستند. در استنباط کلاسیک، مدل‌های زیادی را می‌توان یافت که نیاز به رفع ناشناسایی دارند. در زیر به مثالی از یک مدل کلاسیک در اسوارتز (۲۰۰۴) اشاره می‌کنیم.

مثال ۳.۲.۳. مدل طرح عاملی یک‌طرفه

$$y_{ij} = \mu + \alpha_i + \epsilon_{ij}, \quad i = 1, \dots, k, \quad j = 1, \dots, n_i,$$

^۵Rectifiable

^۶Unidentified or under identified

^۷Just identified

^۸Over identified

را در نظر بگیرید، که در آن، μ میانگین کلی، α_i ها سطوح عاملی و ϵ_{ij} مولفه خطا است که از توزیع $N(0, \sigma^2)$ پیروی می کند. بردار پارامتر $(\mu, \alpha_1, \dots, \alpha_k, \sigma)$ قابل شناسایی نیست زیرا برای همه مقادیر ثابت و معلوم γ ، دو بردار $(\mu, \alpha_1, \dots, \alpha_k, \sigma)$ و $(\mu - \gamma, \alpha_1 - \gamma, \dots, \alpha_k - \gamma, \sigma)$ توزیع احتمال یکسانی از داده ها را نتیجه می دهند.

۳.۳ شناسایی پذیری و اهمیت آن در استنباط بیزی

در سال های اخیر، پیشرفت در ساختارهای نرم افزاری و کامپیوتری و امکان برآزش مدل های مختلف آماری موجب شده است تا مدل های پیچیده با استقبال بیشتری از سوی آماردانان و کاربران آن ها روبه رو شوند. نتیجه این امر آن است که ممکن است برخی مدل ها به قدر کافی درک نشده و ناشناسایی به طور ناخواسته وارد مدل و باعث خسارت شود.

رویارویی استنباط بیزی با مسئله شناسایی پذیری مدل را می توان از دو جنبه فلسفی و کاربردی مورد توجه قرار داد. نگرانی که موجب می شود تا استنباط بیزی از نظر کاربردی در چالش ناشناسا بودن یا نزدیک ناشناسا بودن قرار گیرد، گرایش شدید پارامترهای توزیع پسین به یکدیگر است. در واقع، اهمیت شناسایی پذیری در استنباط بیزی از آن جا آغاز می شود که با حضور مسئله شناسایی پذیری، تمایل به داشتن همبستگی شدید بین پارامترها در توزیع پسین به طور چشمگیری افزایش می یابد و در نتیجه زنجیر MCMC در یک زمان منطقی با همگرایی ضعیف همراه می شود یا به همگرایی نمی رسد (اسوارتز، ۲۰۰۴). نگرانی دوم این است که در صورت وجود مسئله شناسایی پذیری حتی با نمونه با اندازه بزرگ، توانایی غلبه بر اطلاعات پیشین به دست نمی آید (جانسون، ۲۰۰۱). علاوه بر این، مسئله شناسایی پذیری را در مدل های سلسله مراتبی نیز می توان یافت (باسو، ۱۹۸۳). در زیر مثالی از اسوارتز (۲۰۰۴) ذکر شده است که به فهم بیشتر ما از این مسئله کمک می کند.

مثال ۱.۳.۳. فرض کنید $Y = (Y_1, \dots, Y_n)'$ نمونه ای تصادفی از توزیع نرمال با میانگین μ و واریانس ۱ باشد، به طوری که $\mu|\alpha \sim N(\alpha, 1)$ و $\alpha \sim N(\alpha_0, 1)$ ، که در آن α_0 معلوم است. هدف، تعیین توزیع پسین μ و α است. با انجام عملیاتی جبری، می توان نشان داد توزیع توام پسین (μ, α) نرمال دومتغیره با بردار میانگین $(\frac{\sum y_i + \alpha_0}{n+1}, \frac{n\bar{y} + (n+1)\alpha_0}{n+1})'$ و ماتریس کوواریانس با عناصر $\sigma_{11} = \frac{1}{n+1}$ ، $\sigma_{12} = \sigma_{21} = \frac{1}{n+1}$ و $\sigma_{22} = \frac{n+1}{n}$ است. در این جا $\bar{y}$ میانگین مشاهدات است. این مدل سلسله مراتبی شناسایی پذیر نیست، زیرا برای هر $\alpha \neq \alpha'$ ، $F(\mu, \alpha|y) = F(\mu, \alpha'|y)$ ، که در آن y بردار مشاهدات پاسخ و $F(\cdot, \cdot|y)$ در این جا تابع توزیع پسین توام (μ, α) است.

یکی از دلایل اهمیت ناشناسایی این است که همیشه پارامتربندی منحصر به فرد مدل ها برای داشتن قابلیت تفسیرپذیری و برآوردی سازگار، امری ضروری محسوب می شود.

۴.۳ حل مسئله شناسایی پذیری

فرض کنید یک مدل آماری خاص با p پارامتر نامعلوم داریم که به نظر می‌رسد برای داده‌های ما مناسب باشد. اما این مدل قابل شناسایی نیست، به این معنی که مقادیر چندگانه‌ای برای بردار پارامتر با توزیع یکسان برای داده‌های مشاهده‌شده وجود دارند. افراد در طول زمان از روش‌های مختلفی برای حل این مسئله در علم آمار استفاده کرده‌اند. در زیر به برخی از آن‌ها اشاره می‌کنیم.

۱. یکی از راه‌های متداول این مشکل در استنباط بیزی، انتخاب توزیع پیشین است (اسوارتز، ۲۰۰۴). اما این راه حل محکمی نیست، زیرا در این صورت نه تنها داده‌ها بلکه توزیع پیشین نیز توانایی شناسایی پارامترهای مدل را ندارند. از طرفی، افراد در این روش به انتخاب توزیع‌های پیشین خاص محدود می‌شوند.

۲. یک راه حل استاندارد به اعتقاد تحلیل گرانی همچون گستافسون (۲۰۰۴)، استفاده از قیدهای شناسایی مناسب است. با قیدبندی کردن پارامترها، مدل با قدرت شناسایی بالا به دست می‌آید (گاس و همکاران، ۲۰۲۰). به عبارتی، در این حالت می‌توان فضای پارامتر شدنی Ω را به $\bar{\Omega} \subset \Omega$ برای همه $\Theta \subset \Omega$ محدود کرد. در این رساله، هدف ما پیدا کردن قید مناسب برای شناسایی مدل لجیت چندجمله‌ای است.

۳. راه دیگر استفاده از قضیه زیر است.

قضیه ۱.۴.۳. (پائولینو و پیرا، ۱۹۹۴). فرض کنید $\psi(\vartheta)$ تابع شناساگر باشد. تابع $\lambda(\vartheta)$ قابل شناسایی است اگر و تنها اگر تابعی از $\psi(\vartheta)$ باشد.

ما در این رساله، راه حل دوم، قید مجموع مساوی با صفر^۹ را برای رهایی از مشکل شناسایی پذیری در مدل لوجیت چندجمله‌ای پیشنهاد می‌کنیم و نتایج حاصل را با قید متداول رده پایه^{۱۰} یا رده گوشه^{۱۱} مقایسه می‌کنیم.

۱.۴.۳ تاریخچه استفاده از قیدها در آمار

استنباط آماری بر اساس قیدها در طول دهه‌های گذشته مسیر رو به جلو را طی کرده است که می‌توان آغاز آن در سال‌های اولیه را به کودو (۱۹۶۳) و میروین (۱۹۹۴) نسبت داد. بارمی و دیکسترا (۱۹۹۸) روش قید مدل باکس کانوکس بر روی مدل لوجیت چندجمله‌ای را پیشنهاد کردند. لین (۱۹۹۷) روش قید روی عناصر واریانس در مدل خطی تعمیم‌یافته را آزمود. هال و

^۹Sum to zero

^{۱۰}Baseline category

^{۱۱}Corner point category

پریستارت (۲۰۰۱)، قیدهای ترتیبی را بر روی مدل‌های خطی تعمیم یافته و ناخطی تعمیم یافته به کار گرفتند. همچنین بسیاری از افراد روش‌های متفاوت قیدبندی را در شاخه‌های مختلف علم آمار به کار گرفتند که در نتیجه آن دقت پیشگویی در مدل‌ها به طور جدی بالا رفت (سید، ۲۰۲۰). پیشرفت توأم با نیاز روزانه به مدل‌های این چنینی موجب شده است تا افراد مختلفی وارد حوزه استنباط آماری با قیدها شوند. از جمله آن‌ها می‌توان به گاس و همکاران (۲۰۲۰)، سید (۲۰۲۰) و وینالد و همکاران (۲۰۲۱) اشاره کرد.

علی‌رغم پیشرفت عمده در علم داده‌ها و هوش مصنوعی و در نتیجه نیاز به پیش‌بینی دقیق رده چه از منظر پاسخ و چه از منظر احتمال، کمتر کسی را می‌توان یافت که در قیدبندی مدل لوجیت کار کرده باشد. برای مثال، می‌توان به سید (۲۰۲۰) در رهیافت کلاسیک در قید بر روی پارامترهای احتمال مدل لجیت با هدف ارتقاء مدل برای هوش مصنوعی اشاره کرد. ما در این رساله با بهره گرفتن از قیدها، علاوه بر رفع مشکل شناسایی پذیری در مدل لجیت، به این بخش از استنباط آماری نیز توجه کرده‌ایم. اگرچه استفاده از قیدها نیازمند به کارگیری مدل‌های پیچیده‌تر از نظر محاسباتی و زمان اجرا در برنامه‌های کدنویسی است، اما تجربه محققان مختلف در این زمینه نشان داده است که به کار گرفتن قیدها موجب افزایش کارایی و دقت در نتایج به دست آمده شده است (سید، ۲۰۲۰).

در این مرحله، بحث را درباره عملکرد دو نوع پارامتربندی، بر اساس قیدهای مجموع مساوی با صفر و رده گوشه، در مدل لوجیت چندجمله‌ای پی می‌گیریم به گونه‌ای که یک مدل با کیفیت مشابه با مدل لوجیت چندجمله‌ای که مشکل شناسایی پذیری در آن‌ها نیز یکسان باشد را با روش‌های نامبرده مورد بررسی قرار می‌دهیم.

در قید رده پایه، معمولاً افراد پارامترهای مربوط به یکی از رده‌ها (رده پایه) را برابر صفر قرار می‌دهند و در قید مجموع مساوی با صفر، که پیشنهاد ما در این رساله است، حاصل جمع پارامترهای مد نظر برای رده‌ها برابر صفر در نظر گرفته می‌شود. عملکرد این دو نوع قید برای مدل چندجمله‌ای در فصل چهارم به طور کامل تشریح شده است.

۵.۳ مدل نرمال با ویژگی ناشناسایی مشابه با مدل لوجیت چندجمله‌ای

در این بخش برای ارزیابی و یافتن قید مناسب، سعی کردیم قیدهای مد نظر را در یک مدل نرمال که مشکل شناسایی در آن مشابه مدل چندجمله‌ای است، به کار بگیریم. یک توزیع پیوسته مانند توزیع نرمال که در آن پاسخ‌ها به صورت رده‌ای مرتب شده باشند، مشکل شناسایی پذیری در آن مشابه مشکل شناسایی پذیری در مدل لوجیت چندجمله‌ای است (اگرستی، ۲۰۰۲). بر این اساس، سعی کردیم تا با یک مدل شبیه‌سازی شده، عملکرد انواع قیدها را بررسی کرده و در حد امکان نتایج محکمی را ارائه کنیم.

فرض کنید یک توزیع نرمال با میانگین عاملی در سه رده داشته باشیم. فرض کنید Y یک متغیر تصادفی نرمال با میانگین $Z\beta$ است، که در آن Z یک متغیر عاملی با سه سطح، و $\beta = (\beta_0, \beta_1, \beta_2, \beta_3)'$ بردار پارامترهای نامعلوم مدل با مولفه β_0 به عنوان عرض از مبدا است. بنابراین $Y \sim N(Z\beta, Q_y^{-1})$ که در آن، $Q_y = \tau_y I$ ماتریس دقت مدل است. همچنین توزیع نرمال با میانگین صفر و ماتریس دقت $Q_\beta = \tau_\beta I$ را به عنوان توزیع پیشین برای β در نظر می گیریم. هدف ما به دست آوردن توزیع پسین $(\beta_1, \beta_2, \beta_3)'$ است. به همین منظور، برای به دست آوردن توزیع پسین، ابتدا قید رده پایه یا رده گوشه را در نظر می گیریم. برای این کار فرض می کنیم $\eta_{ij} = \beta_0 + \beta_j$ پیشگوی خطی مربوط به میانگین آزمودنی i و رده j باشد، به طوری که در آن $i = 1, \dots, n$ و $j = 1, 2, 3$. حال فرض کنید $\eta_{ij}^* = \beta_0 + \beta_j^*$ ، به طوری که برای اعمال قید، قرار می دهیم $\beta_3^* = 0$. با انجام محاسباتی، می توان نوشت

$$Z_1 \beta^* = \begin{bmatrix} 1 & 1 & 0 \\ 1 & 0 & 1 \\ 1 & 0 & 0 \end{bmatrix} \begin{bmatrix} \beta_0 \\ \beta_1^* \\ \beta_2^* \end{bmatrix}.$$

توزیع $\beta^* \sim N(0, Q_{\beta^*}^{-1})$ را به عنوان توزیع پیشین β^* می پذیریم، که در آن $Q_{\beta^*} = \tau_\beta I_{3 \times 3}$. بنابراین می توان نوشت

$$\begin{aligned} \pi(\beta^* | y) &\propto |Q_y|^{1/2} e^{-\frac{1}{2}(y - Z_1 \beta^*)' Q_y (y - Z_1 \beta^*)} |Q_{\beta^*}|^{1/2} e^{-\frac{1}{2} \beta^{*'} Q_{\beta^*} \beta^*} \\ &\propto e^{-\frac{1}{2} [\beta^* - (Z_1' Q_y Z_1 + Q_{\beta^*})^{-1} Z_1' Q_y y]' (Z_1' Q_y Z_1 + Q_{\beta^*})^{-1} [\beta^* - (Z_1' Q_y Z_1 + Q_{\beta^*})^{-1} Z_1' Q_y y]}. \end{aligned}$$

در نتیجه $\beta^* = (\beta_0, \beta_1^*, \beta_2^*)'$ دارای توزیع پسین نرمال با میانگین $(Z_1' Q_y Z_1 + Q_{\beta^*})^{-1} Z_1' Q_y y$ و ماتریس دقت $(Z_1' Q_y Z_1 + Q_{\beta^*})$ است. با انجام محاسبات ساده، میانگین برای رده ها به صورت

$$\beta_1 = \frac{-1}{3} \beta_1^* + \frac{-1}{3} \beta_2^*, \quad \beta_2 = \frac{2}{3} \beta_1^* + \frac{-1}{3} \beta_2^*, \quad \beta_3 = \frac{-1}{3} \beta_1^* + \frac{2}{3} \beta_2^*.$$

به دست می آید. در گام دوم برای به کار گرفتن قید مجموع مساوی با صفر باید $\sum_{m=1}^3 \beta_m = 0$ را به مدل وارد کنیم. در واقع ما به توزیع پسین β تحت قید خطی $A\beta = 0$ نیاز داریم، که در آن $A = (0, 1, 1, 1)$ و در نتیجه $A\beta = \beta_1 + \beta_2 + \beta_3 = 0$ به شرط $A\beta = 0$ نرمال با میانگین صفر و ماتریس دقت $\tau_\beta^{-1}(I - B)$ است که در آن $B = Q_\beta^{-1} A' (A Q_\beta^{-1} A')^{-1} A$. در ابتدا تعریف می کنیم

$$Z = \begin{bmatrix} 1 & 1 & 0 & 0 \\ 1 & 0 & 1 & 0 \\ 1 & 0 & 0 & 1 \end{bmatrix}.$$

پیشگوی خطی برای $i = 1, \dots, n$ و $j = 1, 2, 3$ برابر است با $\eta_{ij} = \beta_0 + \beta_j$ بنابراین توزیع پسین β به صورت

$$\pi(\beta | y) \propto |Q_y|^{1/2} e^{-\frac{1}{2}(y - Z\beta)' Q_y (y - Z\beta)} |\tau_\beta^{-1}(I - B)|^{1/2} e^{-\frac{1}{2} \beta' (\tau_\beta^{-1}(I - B)) \beta}$$

$$\propto e^{-\frac{1}{2}[\beta - (Z'Q_y Z + \tau_\beta^{-1}(I-B))^{-1}Z'Q_y y]'(Z'Q_y Z + \tau_\beta^{-1}(I-B))[\beta - (Z'Q_y Z + \tau_\beta^{-1}(I-B))^{-1}Z'Q_y y]}$$

نتیجه می‌شود. در نتیجه یک توزیع نرمال با میانگین $(Z'Q_y Z + \tau_\beta^{-1}(I-B))^{-1}Z'Q_y y$ و ماتریس دقت $Z'Q_y Z + \tau_\beta^{-1}(I-B)$ توزیع پسین برای β است، که در آن $(I-B)$ یک ماتریس تکین است (رو و هلد، ۲۰۰۵). بنابراین باید از روش معکوس ماتریس تعمیم‌یافته^{۱۲} و تجزیه مقدار تکین^{۱۳} استفاده کنیم.

در شکل‌های ۱.۳ و ۲.۳ ویژگی‌های توزیع‌های پسین حاصل از دو نوع پارامتر بندی به‌عنوان تابعی از τ_β مقایسه و نمایش داده شده‌اند. شکل ۱.۳ نسبت میانگین‌های پسین را با دقت دو رقم اعشار برای $\tau_\beta \in \{0/001, 0/005, 0/25, 0/5, 0/75, 1\}$ نشان می‌دهد. تفاوتی بین میانگین‌های پسین مشاهده نمی‌شود. شکل ۲.۳ دقت برآورد شده را برای $\tau_y = 1$ و $n = 30$ و مقادیر یکسان τ_β مانند قبل نشان می‌دهد. تفاوت خیلی کمی بین دقت‌های پسینی تنها وقتی که τ_β مقادیر کوچکی دارد، به چشم می‌خورد. در پایان، واریانس حاصل از دو مدل با $\tau_y = 1$ و $\tau_\beta = 0/001$ به‌طوری که اندازه n رو به افزایش است، در شکل ۳.۳ به نمایش گذاشته شده است. همان‌طور که در شکل مشهود است، وقتی اندازه نمونه افزایش یابد، واریانس‌های پسین حاصل از مدل‌ها به یکدیگر همگرا می‌شوند.

با در نظر گرفتن نتایج شبیه‌سازی مدل نرمال عاملی، دو قضیه زیر را بیان و ثابت می‌کنیم. همان‌طور که نتایج شبیه‌سازی نشان دادند، میانگین‌های حاصل از توزیع‌های پسین برای دو قید متفاوت تا دو رقم اعشار با هم برابر هستند که در زیر این برابری را به صورت رسمی نیز اثبات می‌کنیم.

قضیه ۱.۵.۳. مدل نرمال با میانگین عاملی $Z\beta$ و دقت $Q_y = \tau_y I$ را در نظر بگیرید. فرض کنید برای شناسایی مدل از دو قید رده گوشه و مجموع مساوی با صفر استفاده کنیم که ماتریس‌های دقت توزیع‌های پیشین نرمال متناظر آن‌ها به ترتیب Q_{β^*} و Q_β هستند. دو میانگین حاصل از پسین برای دو قید رده گوشه و قید مجموع مساوی با صفر، تقریباً اختلاف ناچیز دارند.

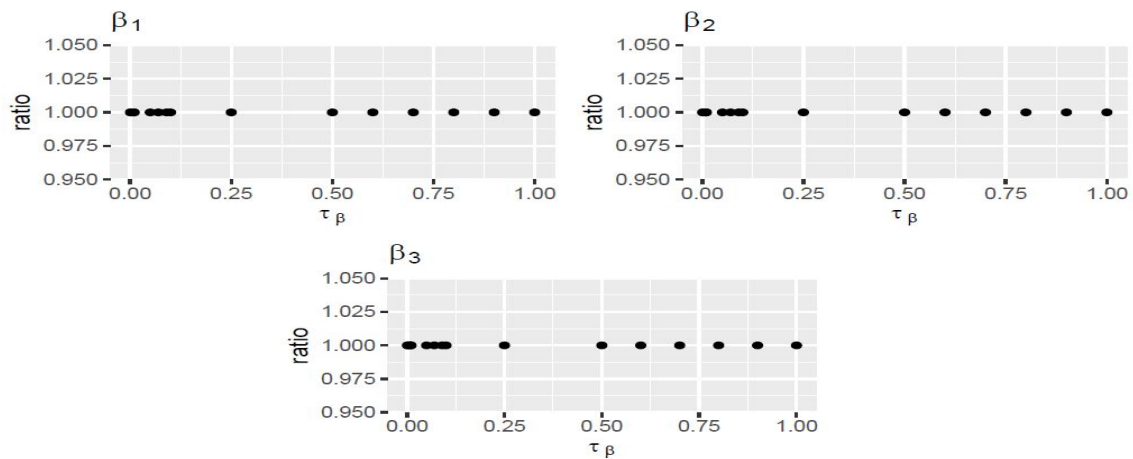
برهان. مدل را بدون عرض از مبدا (برای سبکی فرمول‌ها) در نظر بگیرید. قبل از شروع برهان در ابتدا تعریف می‌کنیم

$$Z = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix},$$

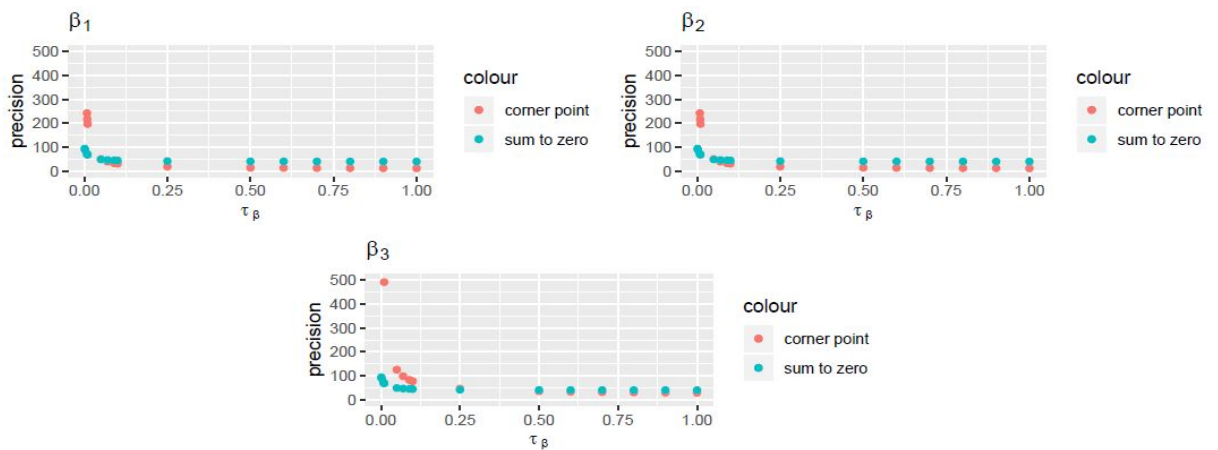
$$Z_1 = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 0 \end{bmatrix}.$$

^{۱۲} Moore-Penrose

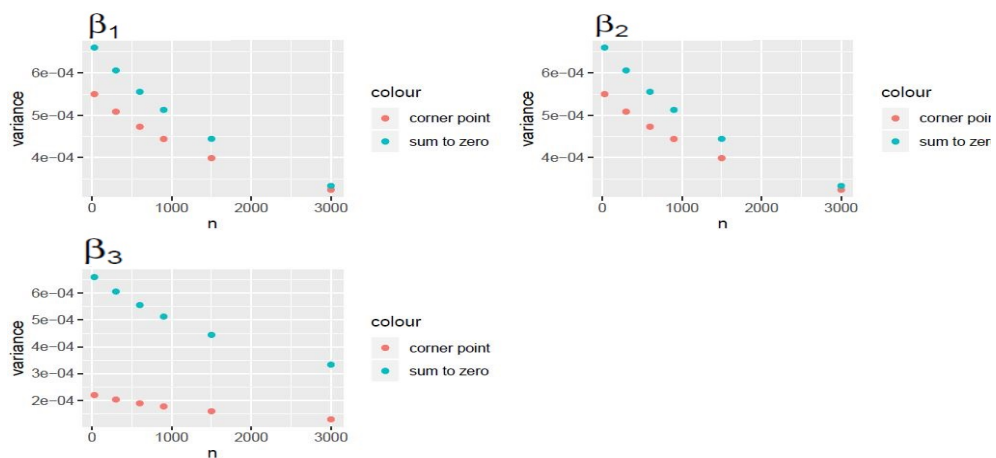
^{۱۳} Singular value dicomposition (svd)



شکل ۱.۳: نسبت میانگین پسین با در نظر گرفتن قیدهای شناسایی پذیری در مدل نرمال عاملی.



شکل ۲.۳: مقایسه دقت با در نظر گرفتن قیدهای شناسایی پذیری در مدل نرمال عاملی برای $n = 300$ و $\tau_y = 1$. دقت‌های برآورد شده در هر رده وقتی τ_β از 100% تا 1% تغییر می‌کند، تقریباً به یکدیگر نزدیک هستند.



شکل ۳.۳: مقایسه واریانس با در نظر گرفتن قیدهای شناسایی پذیری در مدل نرمال عاملی وقتی تعداد مشاهدات رو به افزایش است.

در حالت کلی قرار می‌دهیم $Q_\beta = Q$. می‌دانیم میانگین‌های مدل‌های قید مجموع مساوی با صفر و قید رده گوشه به ترتیب برابر است با:

$$(Z'Q_y Z + Q^{-1}(I - B))^{-1} Z'Q_y y,$$

و

$$(Z'Q_y Z_\lambda + Q_{\beta*})^{-1} Z'Q_y y.$$

با مقایسه دو میانگین در می‌یابیم که تفاوت اصلی بین دو میانگین در ماتریس دقت است. برای نشان دادن میزان این تفاوت نسبت دترمینان ماتریس‌های دقت را به صورت

$$\frac{|Z'Q_y Z + Q^{-1}(I - B)|}{|Z'Q_y Z_\lambda + Q_{\beta*}|} = \frac{|Z'Q_y Z(I + Q^{-1} Z'Q_y^{-1} Z(I - B))|}{|Z'Q_y Z_\lambda + Q_{\beta*}|}$$

می‌نویسیم. داریم

$$\frac{|Z'Q_y Z(Q + Z'Q_y^{-1} Z - Z'Q_y^{-1} ZB)|}{|Q_{\beta*} + Z'Q_y Z_\lambda|}$$

که برابر با

$$\frac{|Z'Q_y Z(Q + Z'Q_y^{-1} Z)(I - (Q + Z'Q_y^{-1} Z)^{-1} Z'Q_y^{-1} ZB)|}{|Q_{\beta*} + Z'Q_y Z_\lambda|},$$

است. با توجه به صورت کسر باید نشان دهیم $1/\tau_y B$ نسبت به $Q + 1/\tau_y I$ از مقدار کوچکی برخوردار است. به عبارتی پیرانتز دوم در صورت کسر تقریباً مقدار صفر می‌گیرد. حال تجزیه ویژه^{۱۴} $Q = V\Lambda V' + \epsilon NN'$ را در نظر بگیرید که در آن $N \in \mathbb{R}^{n \times w}$ و $V \in \mathbb{R}^{n \times (n-w)}$ به گونه‌ای که w تعداد قیدها و n بعد ماتریس دقت است. برای $R \in \mathbb{R}^{w \times w}$ می‌توان نوشت $A = RN'$ و در نتیجه بنابر رو و هلد (۲۰۰۵) داریم: $V'A' = V'NA' = 0$ و همچنین $V'NN' = 0$.

$$Q^{-1}A' = V\Lambda^{-1}V'A' + \epsilon^{-1}NN'NR' = \epsilon^{-1}A'.$$

بسطی از B را به صورت

$$\begin{aligned} B &= Q^{-1}A'(AQ^{-1}A')A \\ &= \epsilon^{-1}NR'(\epsilon^{-1}RN'NR')'RN' \\ &= NN', \end{aligned} \tag{۱.۳}$$

می‌توان نوشت. بنابراین

$$\begin{aligned} (Q + Z'Q_y^{-1} Z)^{-1} Z'Q_y^{-1} ZB &= 1/\tau_y (Q + 1/\tau_y I)^{-1} B \\ &= 1/\tau_y V(\Lambda + 1/\tau_y I)^{-1} V'NN' + 1/\tau_y (\epsilon + 1/\tau_y)^{-1} NN' \\ &= (1 - \frac{\epsilon}{1/\tau_y + \epsilon}) NN'. \end{aligned} \tag{۲.۳}$$

^{۱۴} eigendecomposition

پس

$$\begin{aligned} (I - (Q + Z'Q_y^{-1}Z)^{-1}Z'Q_y^{-1}ZB) &= (I - 1/\tau_y(Q - 1/\tau_y I)^{-1}B) \quad (3.3) \\ &= \left(\frac{\epsilon}{1/\tau_y + \epsilon}\right)^w. \end{aligned}$$

توجه داشته باشید که (۳.۳) به ابرپارامترهای نامعلوم بستگی دارد. اما در مدل، $\epsilon \ll \tau_y$ و بنابراین (۳.۳) تقریباً برابر صفر است. اختلاف در میانگین با اختلاف ناچیز بین دو ماتریس دقت Q و Q_{β^*} کنترل می‌شود. $\square$

در ادامه قصد داریم نشان دهیم در این حالت دقت مدل با قید مجموع مساوی با صفر از قید رده گوشه بیشتر است.

قضیه ۲.۵.۳. مدل بیان شده در قضیه ۱.۵.۳ را در نظر بگیرید. در این حالت دقت مدل حاصل از قید مجموع مساوی با صفر از مدل حاصل از قید رده گوشه بیشتر است.

برهان. می‌دانیم قید رده گوشه دقت برآورد مدل را نسبت به مدل بدون قید کاهش می‌دهد. بنابراین کافی است نشان دهیم که دقت مدل با قید مجموع مساوی با صفر از مدل بدون قید بیشتر است ابتدا تابع درستنمایی مربوط به مدل بدون قید را به صورت

$$\begin{aligned} \pi(\beta|y) &\propto |Q_y|^{1/2} e^{-\frac{1}{2}(y-Z\beta)'Q_y(y-Z\beta)} |Q|^{1/2} e^{-\frac{1}{2}\beta'Q\beta} \\ &\propto e^{-\frac{1}{2}[\beta-(Z'Q_yZ+Q)^{-1}Z'Q_yy]'(Z'Q_yZ+Q)^{-1}[\beta-(Z'Q_yZ+Q)^{-1}Z'Q_yy]}. \end{aligned}$$

داریم. از برهان قضیه قبل به خاطر داریم $I = VV' + NN'$ و $Q = V' + \epsilon NN'$ ، $B = NN'$ معکوس مور-پنروس $(I - B)Q^{-1}$ را با $[(I - B)Q^{-1}]^+$ نشان می‌دهیم. بنابراین داریم

$$[(I - B)Q^{-1}]^+ + Z'Q_yZ)^{-1} = V(\Lambda + \tau_y I)^{-1}V' + \tau_y^{-1}NN'$$

و در نتیجه اختلاف دو دقت برابر است با

$$[(I - B)Q^{-1}]^+ + Z'Q_yZ)^{-1} - (Q + Z'Q_yZ)^{-1} = \frac{\epsilon}{\tau_y(\tau_y + \epsilon)}NN' > 0.$$

که نشان می‌دهد دقت مدل بدون قید از دقت مدل با قید مجموع مساوی با صفر کمتر است. از طرفی می‌دانیم دقت مدل با قید رده گوشه از مدل بدون قید کمتر است. بنابراین دقت مدل با قید مجموع مساوی با صفر از قید رده گوشه بیشتر است. $\square$

قضیه دوم نشان می‌دهد قید مجموع مساوی با صفر یکی از اهداف ما از پیشنهاد این قید که همان افزایش دقت برآورد است را به خوبی ادا کرده است. لازم به ذکر است که این قضایا در حالت کلی ثابت شده‌اند و به هر مدل LGM دیگری نیز قابل تعمیم هستند. بنابراین می‌توان این نتایج را به هر مدلی که بتوان آن را به صورت یک LGM در آورد، نسبت داد.

فصل ۴

مدل چندجمله‌ای جمعی - ساختاری با در نظر گرفتن قید شناسایی

در این فصل، مدل چندجمله‌ای بیزی با ساختار رگرسیونی جمعی را برای داده‌های رسته‌ای فضایی معرفی و عملکرد آن را ارزیابی می‌کنیم. با هدف برازش سریع یک مدل با کیفیت بالا، از تقریب INLA استفاده می‌کنیم. برای این کار، در ابتدا با کمک تبدیل چندجمله‌ای - پواسن (بیکر، ۱۹۹۴) به مدل چندجمله‌ای قابلیت برازش و استنباط با روش INLA می‌بخشیم. برای رفع مشکل شناسایی ذاتی در پارامترهای مدل از دو نوع قید پیشنهادی در فصل قبل استفاده کرده و کارایی آن را به کمک مثال‌های شبیه‌سازی و واقعی ارزیابی می‌کنیم. مثال واقعی شامل داده‌های مربوط به شدت تصادفات رانندگی در استان مازندران است. تاکنون مدل‌های لوجیت چندجمله‌ای به‌طور وسیعی برای تحلیل داده‌های رسته‌ای مربوط به تصادفات جاده‌ای به کار گرفته شده‌اند. برای داشتن یک مقایسه جامع با مدل‌های به کار گرفته‌شده در این زمینه، یک نسخه بیزی از مدل چندجمله‌ای کسری^۱ (تیل، ۱۹۶۰، مولا‌هی، ۲۰۱۰، لی، ۲۰۱۸) را توسعه داده و نتایج حاصل از برازش مدل پیشنهادی خود را با این مدل در مثال واقعی مقایسه و نتیجه‌گیری می‌کنیم.

^۱Fractional split multinomial model

۱.۴ تعریف مدل

در این بخش مدل چندجمله‌ای بیزی را در قالب یک رگرسیون جمعی - ساختاری با مولفه فضایی، با در نظر گرفتن دو نوع قید شناسایی، معرفی می‌کنیم.

۱.۱.۴ رگرسیون جمعی - ساختاری برای داده‌های رسته‌ای اسمی

فرض کنید در هر آزمایش $n, \dots, 1, i$ ، یک انتخاب از بین J رده داشته باشیم. بردار مشاهدات برای هر i به صورت $y_i = (y_{i1}, \dots, y_{iJ})'$ است، به گونه‌ای که $\sum_{j=1}^J y_{ij} = n_i$. همچنین فرض کنید $\pi_i = (\pi_{i1}, \dots, \pi_{iJ})'$ که در آن احتمال انتخاب j امین رده در i امین آزمایش است. درستنمایی مدل چندجمله‌ای برای هر آزمایش i به صورت

$$p(y_i | n_i, \pi_i) = \binom{n_i}{y_i} \pi_{i1}^{y_{i1}} \dots \pi_{iJ}^{y_{iJ}}, \quad (1.4)$$

است، که در آن برای هر $i = 1, \dots, n$ ، $\sum_j \pi_{ij} = 1$ می‌دانیم به کمک مدل لجیت چندجمله‌ای، می‌توان متغیرهای تبیینی را به صورت زیر وارد مدل کرد

$$\pi_{ij} = \frac{\exp(\eta_{ij})}{\sum_{s=1}^J \exp(\eta_{is})} \quad (2.4)$$

که در آن η_{ij} پیشگوی خطی و به صورت

$$\eta_{ij} = \alpha_j + x_{i1j}\beta_{1j} + \dots + x_{iJj}\beta_{Jj}, \quad i = 1, \dots, n, \quad j = 1, \dots, J, \quad (3.4)$$

است. برای j امین رده، همان عرض از مبدا است و $(x_{i1j}, \dots, x_{iJj})'$ برداری از متغیرهای تبیینی برای همه آزمایش‌ها و $\beta_j = (\beta_{1j}, \dots, \beta_{Jj})'$ بردار ضرایب رگرسیونی است. احتمال انتخاب یک گزینه از بین چند گزینه تحت تاثیر سایر گزینه‌ها نیست (لوکه، ۱۹۵۹). در اصطلاح ریاضی، این به آن معنی است که احتمال انتخاب گزینه c از تعداد J گزینه برابر با

$$\pi_c = \frac{w_c}{\sum_j w_j},$$

است، که در آن w_c وزن گزینه c است و به ماکسیمم نرم^۲ معروف است (لوکه، ۱۹۵۹). این نمایش از احتمال رابطه خیلی نزدیکی با مفهوم مدل لجیت در آمار دارد (لوکه، ۱۹۷۷). ما از این حقیقت استفاده می‌کنیم تا قید $\sum_j \pi_{ij} = 1$ را وارد مدل کنیم. در واقع $\sum_{s=1}^J \exp(\eta_{is})$ در رابطه (۲.۴) که به قانون انتخاب لوکه معروف است (لوکه، ۱۹۵۹)، تضمین می‌کند تا مجموع احتمالات در همه رده‌ها یک باشد. به عبارتی دیگر با استفاده از مدل لجیت توانستیم قید شناسایی‌پذیری برای احتمال را وارد مدل کنیم. استفاده از تابع نمایی در تابع لجیت برای

^۲Soft max

سادگی ریاضی و راحتی کار است. یک ویژگی خوب این تابع این است که مشتق آن با خود تابع یکسان است.

با در نظر گرفتن $u = (u_1, \dots, u_{n_f})'$ ، برای نمایش اطلاعات مربوط به اثرات متغیرهای ساختاری می‌توان پیشگوی خطی را به صورت

$$\eta_{ij} = \alpha_j + \sum_{r=1}^{n_f} f_j(u_{ri}) + \sum_{k=1}^l x_{kij} \beta_{kj}, \quad (4.4)$$

توسعه داد، که در آن $f(\cdot)$ نشان‌دهنده تابعی ناخطی از متغیرهای ساختاری (تبیینی پیوسته) است که می‌تواند به صورت اثرات ناپارامتری، زمانی، فضایی و فضایی-زمانی ظاهر شود. اگر به پارامترهای $\alpha = (\alpha_1, \dots, \alpha_J)'$ ، $\beta = (\beta'_1, \dots, \beta'_J)'$ و $\{f(\cdot)\}$ توزیع پیشین گوسی نسبت دهیم، آن‌گاه با یک مدل LGM روبرو هستیم.

۲.۱.۴ مدل‌بندی ساختار وابستگی فضایی

بسته به پیوسته (چگال) یا گسسته (ناچگال) بودن ناحیه تحت مطالعه، مدل‌های متفاوتی برای مدل‌بندی ساختار وابستگی فضایی معرفی شده‌اند. مدل‌های معروف پیشنهادی برای ناحیه اندیس‌گذاری شده گسسته، برای استفاده در INLA توصیه و همچنین در بسته نرم‌افزاری R-INLA تعبیه شده‌اند؛ کاربردهای متعددی از آن‌ها را می‌توان در اسکرود و همکاران (۲۰۱۲)، اگرته و همکاران (۲۰۱۴) و بلانگیاردو و کاملتی (۲۰۱۵) مشاهده کرد.

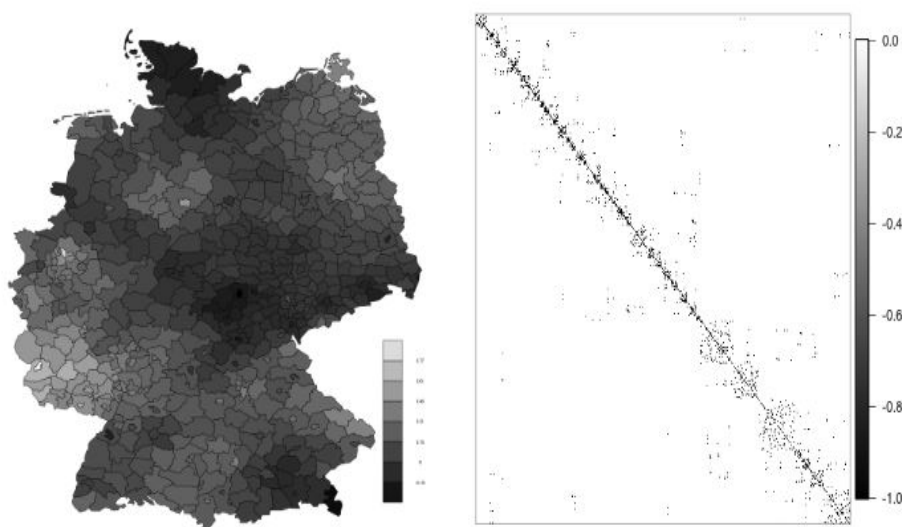
یکی از معروفترین مدل‌های فضایی ناحیه گسسته، مدل بی‌سیگ (بی‌سیگ، ۱۹۹۱) یا همان ICAR است که به افتخار آماردانی با همین نام، ژولین بی‌سیگ، نامگذاری شده است. ما در این رساله از عبارت CAR برای هر مدل گوسی چندمتغیره که از توزیع‌های شرطی ساخته شده باشد، استفاده می‌کنیم. بنابراین، در این رساله مدل بی‌سیگ، یک مثال خاص از مدل CAR محسوب می‌شود.

مدل بی‌سیگ

صورت کلی مدل ICAR را در فصل اول معرفی کردیم. این مدل، اثر فضایی در هر ناحیه، u_i را روی یک مجموعه از نواحی که با $\{1, \dots, n\}$ اندیس‌گذاری و روی نواحی همسایه شرطی شده است، مدل‌بندی می‌کند. معمولاً دو ناحیه را همسایه در نظر می‌گیرند، اگر مرز مشترکی داشته باشند (شکل ۱.۴ را ببینید). توزیع شرطی برای u_i به صورت

$$u_i | \mathbf{u}_{-i}, \tau_u \sim N\left(\frac{1}{d_i} \sum_{j \sim i} u_j, \frac{1}{d_i \tau_u}\right), \quad (5.4)$$

است، که در آن $i \sim j$ همسایه بودن دو ناحیه i و j را نشان می‌دهد و d_i تعداد همسایه‌های ناحیه i است. همچنین τ_u پارامتر دقت مدل است. در این صورت، توزیع توأم $\mathbf{u} = (u_1, \dots, u_n)'$ به



شکل ۱.۴: مدل بی‌سیگ. شکل سمت چپ مربوط به داده‌های یک مثال از نواحی کشور آلمان است. شکل سمت راست ماتریس دقت تنک Q متناظر با این نواحی را نشان می‌دهد.

صورت

$$u|\tau_u \sim N\left(\circ, \frac{1}{\tau_u} Q^{-1}\right),$$

است، که در آن Q یک ماتریس دقت ساختاری است و مولفه‌های آن عبارتند از

$$Q_{i,j} = \begin{cases} d_i, & i = j \\ -1, & i \sim j \\ \circ, & \text{سایر موارد} \end{cases} \quad (۶.۴)$$

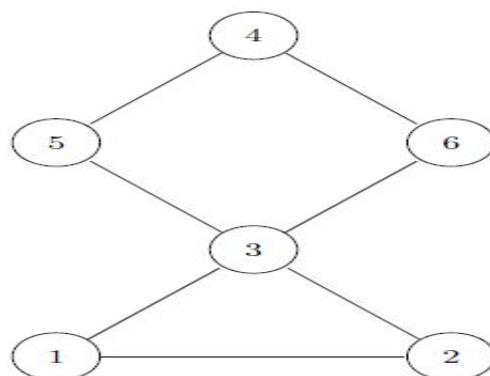
این ماتریس دقت ساختاری، به‌طور مستقیم ساختار همسایگی نواحی را لحاظ می‌کند و با داشتن $O(n)$ عضو ناصفر، تنک است (پنل راست شکل ۱.۴ را ببینید). توجه داشته باشید مدل بی‌سیگی که در آن وزن‌های مثبت برای هر جفت از ناحیه‌های همسایه منظور شده باشد، به‌طور مشابه قابل تعریف است (رو و هلد، ۲۰۰۵). راه‌های مختلفی برای تعیین مولفه‌های ماتریس دقت در (۶.۴) وجود دارند که به‌عنوان نمونه می‌توان به دو روش زیر اشاره کرد.

۱. از طریق ماتریس همسایگی متقارن با بعد $n \times n$

۲. با استخراج مستقیم ساختار از فایل شیپ^۳ به کمک بسته‌های نرم‌افزاری مانند بسته‌های maptools (بیوند و لوین کج، ۲۰۱۷) و spdep (بیوند و پیراس، ۲۰۱۵؛ بیوند و همکاران، ۲۰۱۳) در محیط R.

^۳Shape file

تعریف مدل بی‌سیگ با اختصاص یک گراف شروع می‌شود. گراف از مجموعه‌ای از نقاط و یال، که نواحی و همسایگی بین آن‌ها را نشان می‌دهند، تشکیل می‌شود. یک گراف متصل شده است اگر مسیر وجود داشته باشد (یعنی مجموعه‌ای از یال‌های به هم پیوسته داشته باشد) که هر نقطه را حداقل به یک نقطه دیگر متصل می‌کند. در یک گراف متصل شده، رابطه همسایگی واضح است (هر نقطه حداقل یک همسایه دارد). در شکل ۲.۴ نمونه‌ای از یک گراف (فرنی-استرینو، ۲۰۱۸) را می‌توانید ببینید. مدل بی‌سیگ از چنین گرافی استفاده می‌کند.



شکل ۲.۴: نمونه‌ای از یک گراف متصل شده.

مشکل شناسایی‌پذیری در مدل بی‌سیگ

مدل معرفی‌شده در (۵.۴) و (۶.۴)، ذاتی (ناسره) است به این معنی که میانگین کلی در آن نامشخص رها شده و نیاز به شناسایی دارد. در حالت کلی، دو راه برای شناسایی آن وجود دارد:

۱. **استفاده از قید خطی:** فرنی - اسرینو در سال ۲۰۱۸ نشان داد که با استفاده از یک قید خطی می‌توان مدل را قابل شناسایی کرد. با توجه به قابلیت INLA در به کار گرفتن قید خطی، استفاده از قید مجموع مساوی با صفر را پیشنهاد کرده و به کار می‌گیریم، به طوری که مجموع اثرات فضایی برابر صفر شود، یعنی $\sum_{i=1}^n u_i = 0$.

۲. **اضافه کردن یک مقدار خیلی کوچک به قطر اصلی ماتریس دقت:** در این حالت با افزودن یک مقدار کوچک $\gamma > 0$ به تمام اعضای قطر اصلی ماتریس دقت Q آن را به یک ماتریس معین مثبت تبدیل می‌کنند. بنابراین کافی است نشان دهیم که اختلاف این تقریب با Q تابعی از γ است؛ به عبارتی این اختلاف با γ کنترل می‌شود. زیرا با کمی کندوکاو بیشتر دریافتیم که افزودن قید مجموع مساوی با صفر به مدل به اندازه γ به قطر اصلی ماتریس دقت اضافه می‌کند. بنابراین به سراغ روشی رفتیم که در آن به اختیار خودمان به اندازه γ به قطر اصلی اضافه کردیم و کارایی آن را سنجیدیم.

قضیه ۱.۱.۴. فرض کنید X یک میدان تصادفی گوسی مارکوفی با ماتریس دقت Q با بعد n باشد که ذاتی هم هست. تفاوت بین نرم قوی^۴ و نرم ضعیف^۵ دو ماتریس دقت برابر است با $c\gamma$ که در آن c یک مقدار ثابت است.

لازم به ذکر است که $X|X$ معادل با افزودن قید مجموع مساوی با صفر به مدل است که به اندازه γ به مدل اضافه می‌کند. برای اثبات قضیه، در ابتدا نرم قوی و نرم ضعیف را از رو و هلد (۲۰۰۵) تعریف می‌کنیم.

تعریف ۱.۱.۴. فرض کنید M یک ماتریس $n \times n$ است. نرم قوی M برابر است با

$$\|M\|_s = \max_{x: x'x=1} (x'M'Mx)^{1/2}.$$

همچنین، نرم ضعیف M برابر است با

$$\|M\|_w = \left(\frac{1}{n} \sum_{ij} M_{ij}^2 \right)^{1/2}.$$

برهان. فرض کنید $\{(\lambda_i, e_i), i = 1, \dots, n\}$ جفت‌های مقادیر ویژه و بردار ویژه ماتریس Q هستند. بنابراین رابطه

$$Q = \sum_i \lambda_i e_i e_i' = v \Lambda v',$$

برقرار است، که در آن $\lambda_1 = 0$ ، $\Lambda = \text{diag}(\lambda_1, \dots, \lambda_n)$ ، $vv' = 1$ ، $v' = (e_1, \dots, e_n)$ و $e_1 = 1$. توجه داشته باشید که بردارهای ویژه یکسان باقی می‌مانند. جفت‌های مقادیر ویژه و بردار ویژه ماتریس شرطی X به شرط $1'X$ برابرند با

$$(0, 1), (\lambda_2 + \gamma, e_2), \dots, (\lambda_n + \gamma, e_n)$$

زیرا با استفاده از رو و هلد (۲۰۰۵)

$$\text{Prec}(X|X) = v \tilde{\Lambda} v'$$

که در آن $\tilde{\Lambda} = \text{diag}(0, \lambda_2 + \gamma, \dots, \lambda_n + \gamma)$ و $\text{Prec}(X)$ ماتریس دقت X است. بنابراین نتیجه

$$\|\text{Prec}(X|X) - Q\| = c\gamma,$$

حاصل می‌شود، که در آن مقدار c برای نرم قوی برابر با یک و در غیر این صورت برابر با $\sqrt{\frac{n-1}{n}}$ است. $\square$

^۴ Strong norm

^۵ Weak norm

یادآور می‌شویم که نرم قوی و نرم ضعیف را می‌توان با مقادیر ویژه نیز حساب کرد به‌طوری که برای یک ماتریس دقت با مقادیر ویژه $\lambda_1, \dots, \lambda_n$ نرم قوی و نرم ضعیف به ترتیب برابر است با

$$\|M\|_s = \max_k \lambda_k, \quad k = 1, \dots, n$$

و

$$\|M\|_w = \frac{1}{n} \text{trace}(M'M) = \frac{1}{n} \sum_k \lambda_k, \quad k = 1, \dots, n.$$

با کمی تحقیق بیشتر دریافتیم اگر کسی توزیع پیشین با دقت $\bar{Q} = Q + \gamma I$ را برای مدل فضایی در نظر بگیرد، اگرچه اضافه کردن این مقدار به قطر ماتریس، مدل را سره^۶ می‌سازد، اما فضای صفر^۷ در ماتریس دقت همچنان باقی است. در نتیجه، راه حل اول محکم‌تر و مفیدتر در برطرف شدن مشکل به‌نظر می‌رسد.

در این رساله، برخلاف زمان زیادی که صرف پیدا کردن و برهان قضیه در راه حل دوم شد، از روش اول برای رسیدن به هدف استفاده می‌کنیم چرا که سعی کردیم تمرکز خود را در استفاده از یک روش با اجرای ساده‌تر و در عین حال کاراتر بالا ببریم.

مقیاس‌بندی مدل بی‌سیگ

در این بخش درباره یک عقیده کلیدی صحبت می‌کنیم که به‌کار گرفتن آن موجب می‌شود تا مدل‌هایی که در واقع یک میدان تصادفی گوسی ذاتی^۸ هستند، از نظر تفسیر پارامتر راحت‌تر و قابل مقایسه با سایر مدل‌ها شوند.

تا این جا دریافتیم که در مدل رگرسیونی سلسله‌مراتبی بیزی، یک روش متداول در مدل‌بندی فضایی استفاده از مدل‌های IGMRF است. دقت کنید که منظور از IGMRF همان مدل ICAR است که میدان تصادفی آن گوسی فرض می‌شود. بنابراین، در این رساله که فقط از پذیره گوسی بودن میدان تصادفی استفاده می‌کنیم، این دو عبارت را به‌جای هم به کار می‌بریم. در مدل IGMRF یک ماتریس دقت مقیاس‌بندی‌شده وجود دارد که ساختار همسایگی مدل را منعکس می‌کند و پارامتر مقیاس در آن به عنوان پارامتر دقت تصادفی گزارش می‌شود. پیشین انتخاب‌شده برای پارامتر دقت روی درجه هموارسازی تاثیر می‌گذارد و در نتیجه تأثیر قوی بر روی نتایج پسین دارد (کرشیانکی، ۲۰۱۸). در واقع هر IGMRF یک پارامتر دقت دارد که تفاوت‌های بین همسایه‌ها را کنترل می‌کند. مقادیر بزرگ نشان از رابطه‌ای قوی (همواری بیشتر) دارد، در حالی که مقادیر کوچک‌تر نشان از رابطه ضعیف (همواری کمتر) دارد. برای

^۶ Proper

^۷ Null space

^۸ Intrinsic Gaussian Markov Random Fields (IGMRF)

مثال، برناردلی و همکاران (۱۹۹۵) را ببینید. بنابراین، یکی از چالش‌های پیش رو این هست که τ_u تفسیرپذیر نیست زیرا به گراف تحت مطالعه بستگی دارد. اولین مرحله برای بخشیدن قابلیت تفسیرپذیری به τ_u این هست که مدل بی‌سیگ را از ذاتی بودن در بیاوریم (که در بخش قبل به شرح کامل آن پرداختیم). دومین مرحله این است که پارامتر τ_u نسبت به تمام ناحیه‌ها و مجموعه یال‌ها، پایدار^۹ شود. برای این که ناپایداری را دریابیم، تعریف انحراف استاندارد کناری

$$sd(u_i) = \frac{\sqrt{Q_{ii}^-}}{\sqrt{\tau_u}}$$

را ارایه می‌کنیم، که در آن Q^- واریانس تعمیم‌یافته Q است. انحراف استاندارد کناری روی i (یعنی روی نواحی) ثابت نیست و ماکسیمم و هر نوع میانگین آن‌ها به شدت به n وابسته است. بنابراین، سربای و رو (۲۰۱۴) استفاده از توزیع پیشین بر اساس نوع IGMRF را پیشنهاد کردند. در روش پیشنهادی آن‌ها، ماتریس ساختار همسایگی که به نوع گراف وابسته است، به‌گونه‌ای مقیاس‌بندی می‌شود که میانگین هندسی مقادیر واریانس‌های کناری برابر با یک شود. به این ترتیب، مدل بی‌سیگ با مدل‌های دیگر قابل مقایسه می‌شود. این کار با نقش بستن پارامتر دقت تصادفی بر انحراف استاندارد کناری مدل IGMRF شکل می‌گیرد و در نهایت توزیع پیشین روی τ_u مجدد محاسبه می‌شود. برای توضیح بیشتر، فرض کنید اثر تصادفی u یک IGMRF با میانگین صفر و ماتریس دقت $\tau_u R$ باشد. ماتریس دقت مقیاس‌بندی‌شده را با $\tau_u R \cdot \sigma_{GV}^2(R)$ نشان می‌دهیم که $\sigma_{GV}^2(R)$ واریانس تعمیم‌یافته R است که توسط سربای و رو (۲۰۱۴) به صورت

$$\sigma_{GV}^2(R) = \exp\left(\frac{1}{n} \sum_{i=1}^n \log(\text{diag}(R^-))\right),$$

پیشنهاد شده است. در این جا R^- ماتریس تعمیم‌یافته R است. با انجام این کار واریانس کناری مربوط به هر یک از عنصرهای بردار تصادفی u برابر یک می‌شود.

۳.۱.۴ شناسایی‌پذیری در مدل‌های چندجمله‌ای

از آن جایی که قید $\sum_j \pi_{ij} = 1$ در مدل پیشنهادی موجب می‌شود تا پارامترهای مدل منحصر به فرد نباشند، بنابراین پارامترهای مدل قابل شناسایی نیستند. برای مثال در برآورد β_1 در معادله (۳.۴)، اگر هر مقدار ثابتی را به همه J پارامتر $(\beta_1, \dots, \beta_J)$ اضافه کنیم، تابع درستنمایی یکسان باقی می‌ماند، به‌طوری که داده‌ها توانایی شناسایی پارامترها را ندارند (جین و همکاران، ۲۰۱۳). در این جا، ما به کمک دو قید رده گوشه و مجموع مساوی با صفر مشکل ناشناسا بودن مدل را مرتفع می‌کنیم.

^۹Stable

قید رده گوشه برای مدل چندجمله‌ای

یک از متداول‌ترین روش‌ها برای برطرف کردن مشکل شناسایی‌پذیری در مدل چندجمله‌ای، استفاده از قید رده گوشه است که در آن پارامترهای مربوط به یکی از رده‌ها را برابر با صفر در نظر می‌گیرند. در نتیجه استفاده از این روش، افراد در عین روبرو بودن با J رده قادرند تنها پارامترهای مربوط به $J - 1$ رده را برآورد کنند. برای مثال مدل (۳.۴) را در نظر بگیرید. اگر رده J ، رده گوشه باشد، بنابراین برای همه i ها، $\eta_{iJ} = 0$ یا β_{kJ} ها و α_J مساوی با صفر است. بنابراین بقیه پارامترها باید نسبت به رده J تفسیر شوند. به کار گرفتن این روش در استنباط بیزی به این معنی است که توزیع پیشین مربوط به رده گوشه با رده‌های دیگر متفاوت است. بنابراین، پیشین قابلیت تغییرپذیری ندارد. در نتیجه نتایج مربوط به توزیع پسین تحت تاثیر انتخاب رده گوشه است (کانگدان، ۲۰۰۵). مشکل مشابهی را می‌توان در توزیع‌های با پاسخ عاملی نیز یافت (اگرستی، ۲۰۱۳) که در فصل سوم به شرح مفصل آن پرداختیم. لازم به ذکر است اگر یک رده با مشاهدات کمتر به عنوان رده پایه انتخاب شود، نتایج حاصل پایدار نیست (سید، ۲۰۲۰).

قید مجموع مساوی با صفر برای مدل چندجمله‌ای

قید مجموع مساوی با صفر از نظر مفهومی مشابه مرکزی کردن متغیر پیوسته است، که در آن میانگین را از متغیر پیشگو کم می‌کنند. برای مثال، در مدل (۳.۴) قید مجموع مساوی با صفر به صورت $\sum_j \alpha_j = 0$ و $\sum_j \beta_{kj} = 0$ خواهد بود. در این حالت، رده‌ها تغییرپذیر هستند. در این فصل از رساله به تحلیل و مقایسه قید پیشنهادی با حالت متداول یعنی قید رده گوشه می‌پردازیم.

مدل‌های بیزی سلسله‌مراتبی یکی از ابزارهای مهم روز در مدل‌بندی با ساختارهای پیچیده در فرآیندهای فضایی محسوب می‌شوند. در این حالت، معمولاً تابع درستنمایی مدل صورت بسته‌ای ندارد و آماردانان از روش‌های شبیه‌سازی مونت کارلویی برای تقریب توزیع پسین استفاده می‌کنند. اما همان‌گونه که پیش‌تر اشاره شد، متأسفانه این روش از نظر محاسباتی گران است (ارمرود و وند، ۲۰۱۰؛ مینانسی و همکاران، ۲۰۱۱؛ میلدج و همکاران، ۲۰۱۲؛ بلی و همکاران، ۲۰۱۷). همچنین، به‌طور ویژه، برای مدل‌های چندجمله‌ای علاوه بر هزینه محاسباتی بالا، حتی ممکن است به همگرایی نرسد (دپراتره و وندبروک، ۲۰۱۷؛ بانسال و همکاران، ۲۰۲۰). در مقابل می‌توان از روش INLA، که در فصل دوم معرفی شد، بهره گرفت. با این وجود INLA تا کنون برای مدل پیشنهادی چندجمله‌ای به کار گرفته نشده است، زیرا در این مدل، متغیر پاسخ به صورت ناخطی به عناصر پیشگوی خطی وابسته است^{۱۰}. برای رهایی از این مشکل از تبدیل چندجمله‌ای-پواسن^{۱۱} که اولین بار توسط بیکر (۱۹۹۴) معرفی شد،

^{۱۰}<https://www.r-inla.org/faq>

^{۱۱}Multinomial-Poisson transformation (MPT)

بهره می‌گیریم تا بتوانیم هسته درست‌نمایی مدل چندجمله‌ای را به هسته یک مدل پواسون تبدیل و از قابلیت روش INLA استفاده کنیم. البته این تبدیل هزینه هم دارد که در مقابل کارایی محاسباتی روش INLA ناچیز است.

۴.۱.۴ تبدیل چندجمله‌ای - پواسن

فرض کنید برای هر آزمایش i ، $i = 1, \dots, n$ ، $j = 1, \dots, J$ ، $\{y_{ij}, i = 1, \dots, n, j = 1, \dots, J\}$ تحقق‌هایی از Y_{ij} ها است که دارای توزیع چندجمله‌ای با پارامترهای احتمال $\{\pi_{ij}, i = 1, \dots, n, j = 1, \dots, J\}$ هستند و توسط مدل (۲.۴) به متغیرهای تبیینی x متصل می‌شوند. تابع درست‌نمایی مدل به صورت

$$L_M = \prod_{i=1}^n \prod_{j=1}^J \left(\frac{\exp(\eta_{ij})}{\sum_{s=1}^J \exp(\eta_{is})} \right)^{y_{ij}},$$

قابل نمایش است. چالش اصلی هنگام کار کردن با L_M به دلیل وجود $\Gamma_i = \sum_{s=1}^J \exp(\eta_{is})$ رخ می‌دهد، زیرا موجب می‌شود تا استفاده از لگاریتم بی‌اثر شود. با استفاده از تبدیل MP، Γ_i به‌عنوان یک پارامتر اضافی قابل برآورد جایگزین می‌شود. تبدیل حاصل به صورت

$$L_P = \prod_{i=1}^n \prod_{j=1}^J \left\{ \exp(\eta_{ij} + \phi_i) \right\}^{y_{ij}} \exp\{-\exp(\eta_{ij} + \phi_i)\},$$

است، که در آن $\phi = (\phi_1, \dots, \phi_n)'$ بردار پارامترهای اضافه است که موجب می‌شود تا تبدیل MP موثر و قابل استفاده شود. تابع درست‌نمایی L_P تابعی از پارامترهای مدل و بردار ϕ است. با استفاده از سرافینی (۲۰۱۸)، برآورد ماکسیمم درست‌نمایی پارامترهای ϕ_i در نسبت با Γ_i تعریف می‌شود، به‌طوری که $\exp(\phi_i) \propto \frac{1}{\Gamma_i}$. بنابراین، نتایج مشاهده‌شده بر وجود رابطه $L_M \propto L_P$ تاکید می‌کند. قبل از تعبیه این نتیجه برای توسعه مدل بیزی چندجمله‌ای، برای تشریح مفید بودن این تبدیل، ابتدا چند مثال که تبدیل MP در آن‌ها مفید است را بیان می‌کنیم.

مثال ۱.۱.۴. مدل لگ خطی. پالمگرن (۱۹۸۱) حالت خاصی از تبدیل چندجمله‌ای - پواسن را برای مدل لگ خطی یافت. فرض کنید y_{ij} فراوانی مشاهده‌شده خانه ij ام از یک جدول توافقی $I \times J$ را نشان دهد. فرض کنید $\{Y_{ij}\}$ دارای توزیع چندجمله‌ای با پارامترهای $\{\frac{\eta_{ij}}{\sum_j \eta_{ij}}\}$ است، که در آن $\eta_{ij} = \exp(\lambda^i + \lambda^j)$ شامل اثرات متغیرهای سطری و ستونی است. هسته تابع درست‌نمایی و تبدیل MP متناظر آن به صورت زیر هستند:

$$L_M = \prod_{i=1}^I \prod_{j=1}^J \left(\frac{\exp(\lambda^i + \lambda^j)}{\sum_j \exp(\lambda^i + \lambda^j)} \right)$$

$$Y_{ij} \sim \text{Poi}(\exp(\phi_i + \lambda^i + \lambda^j)).$$

مثال ۲.۱.۴. وایتد (۱۹۸۰) تبدیل MP را برای یک مدل کاکس (کاکس، ۱۹۷۲) با متغیرهای تبیینی رسته‌ای پیشنهاد کرد. فرض کنید i اندیس مجموعه افراد در معرض خطر باشد که برای

آن $i = 1, \dots, n$ ، و j اندیس مربوط به سطوح متغیر تبیینی باشد، به گونه‌ای که $j = 1, \dots, J$. قرار می‌دهیم $y_{ij} = 1$ اگر فرد i که در مجموعه خطر قرار دارد، دارای متغیر تبیینی سطح j باشد؛ در غیر این صورت $y_{ij} = 0$. فرض کنید N_{ij} نشان‌دهنده افرادی با سطح j ام متغیر تبیینی است که در مجموعه خطر i قرار دارند. در این صورت، هسته تابع درستنمایی به شکل

$$L = \prod_{i=1}^n \prod_{j=1}^J \left(\frac{\exp(x_{ij}\beta)}{\sum_{u=1}^J \exp(N_{ij}(x_{iu}\beta))} \right),$$

است. برای این که بتوانیم از تبدیل MP استفاده کنیم، باید درستنمایی را در مقدار ثابت $\prod_i \prod_j N_{ij}^{y_{ij}}$ ضرب کنیم. لازم به ذکر است این عمل تغییری در برآورد ماکسیمم درستنمایی و واریانس جانبی ایجاد نمی‌کند. بنابراین هسته درستنمایی و تبدیل MP به صورت زیر خواهند بود:

$$L_M = \prod_{i=1}^n \prod_{j=1}^J \left(N_{ij} \frac{\exp(x_{ij}\beta)}{\sum_j \exp(x_{ij}\beta)} \right),$$

$$Y_{ij} \sim \text{Poi}(\exp(\phi_i + x_{ij}\beta)).$$

مثال ۳.۱.۴. مدل راش تعمیم‌یافته: داروچ (۱۹۸۱)، مک کولاح (۱۹۸۲)، کرسی و هلاندر (۱۹۸۳)، کانوی (۱۹۸۹، ۱۹۹۲) و اگرستی (۱۹۹۳) از تبدیل MP برای تعمیم مدل راش^{۱۲} استفاده کردند. مدل راش تعمیم‌یافته برای داده‌های چندرده‌ای، احتمال این که فرد i به سوال k پاسخ j را بدهد، به صورت

$$P(j \text{ | } i \text{ شخص}, k \text{ سوال}) = \frac{\exp(\alpha_{ij} + x_k \beta_j)}{1 + \sum_{j=2}^J \exp(\alpha_{ij} + x_k \beta_j)},$$

در نظر بگیرید (کانوی، ۱۹۸۹). شرطی کردن روی تعداد پاسخ‌ها در هر سطح باعث می‌شود تا پارامترهای مربوط به α حذف شوند. برای مثال، سه سطح از پاسخ را بر روی سه سوال در نظر بگیرید. فرض کنید فرد ۱ به ترتیب سطوح پاسخ ۲، ۳ و ۲ را برای سوال‌های ۱، ۲ و ۳ انتخاب کند. در این صورت، هسته درستنمایی برای فرد ۱ به صورت

$$L^{(1)} = \frac{\exp(\alpha_{12} + x_1 \beta_2)}{1 + \sum_{j=2}^3 \exp(\alpha_{1j} + x_1 \beta_j)} \frac{\exp(\alpha_{13} + x_2 \beta_3)}{1 + \sum_{j=2}^3 \exp(\alpha_{1j} + x_2 \beta_j)} \frac{\exp(\alpha_{12} + x_3 \beta_2)}{1 + \sum_{j=2}^3 \exp(\alpha_{1j} + x_3 \beta_j)}$$

است. زیرا سه مجموعه ممکن از دو پاسخ در ۲ سطح و یک پاسخ در ۳ سطح به ترتیب $\{2, 2, 3\}$ ، $\{2, 3, 2\}$ و $\{3, 2, 2\}$ وجود دارند. بنابراین، هسته درستنمایی شرطی برای فرد ۱ به صورت

$$L_M^{(1)} = \frac{\exp(x_1 \beta_2 + x_2 \beta_3 + x_3 \beta_2)}{\exp(x_1 \beta_2 + x_2 \beta_2 + x_3 \beta_3) + \exp(x_1 \beta_2 + x_2 \beta_3 + x_3 \beta_2) + \exp(x_1 \beta_3 + x_2 \beta_2 + x_3 \beta_2)},$$

است. هسته درستنمایی شرطی برای همه مشاهدات برابر است با

$$L_M = \prod_i L_M^{(i)},$$

^{۱۲} Rasch

که با گروه‌بندی افراد می‌توان آن را به شکل ساده‌تری نوشت. فرض کنید s نشان‌دهنده یک مجموعه به شکل {تعداد پاسخ در سطح ۱، تعداد پاسخ در سطح ۲، تعداد پاسخ در سطح ۳} باشد. همچنین فرض کنید J_i زیرمجموعه‌های (اندیس‌گذاری شده با j) از سطوح پاسخ برای $k-1$ سوال که به مجموعه i مربوط می‌شود را نشان دهد. به عنوان نمونه، در مثال بالا $i = \{0, 2, 1\}$ و $J_{\{0, 2, 1\}} = \{\{2, 2, 3\}, \{2, 3, 2\}, \{3, 2, 2\}\}$. فرض کنید y_{ij} تعداد افرادی را نشان دهد که به زیرمجموعه J_i تعلق دارند. فرض کنید برای $k \in j$ ، $\delta_{kc} = 1$ ، یعنی پاسخ برای k از زیرمجموعه j برابر c است؛ در غیر این صورت $\delta_{kc} = 0$. هسته درست‌نمایی شرطی و تبدیل MP متناظر به ترتیب زیر هستند:

$$L_M = \prod_{i=1}^n \prod_{j \in J_i} \left(\frac{\exp(\sum_{k \in j} \sum_c x_{kc} \beta_c \delta_{kc})}{\exp(\sum_{u \in J_i} (\sum_{k \in j} \sum_c x_{kc} \beta_c \delta_{kc}))} \right)^{y_{ij}}$$

$$Y_{ij} \sim \text{Pois}(\exp(\phi_i + \sum_{k \in j} \sum_c x_{kc} \beta_c \delta_{kc})).$$

۵.۱.۴ تبدیل چندجمله‌ای - پواسن بیزی

بیکر (۱۹۹۴) نتایج کلی را در رابطه با معادل بودن برآورد ماکسیمم درست‌نمایی حاصل از درست‌نمایی پواسن با درست‌نمایی چندجمله‌ای و همچنین ویژگی‌های مجانبی واریانس‌های مربوط به آن‌ها ارائه داد. این نتایج به مدل‌های آماری متعددی قابل تعمیم است و توانست سال‌ها برای محققان سادگی محاسبات را به همراه بیاورد.

عقیده بیکر در سال ۲۰۰۴ توسط افرادی همچون سیمان و ریچاردسون در قالب مدل‌های کنترل پا به عرصه مدل‌بندی بیزی گذاشت. هدف آن‌ها رسیدن به نتایجی مشابه با نتایج بیکر اما در پیکره بیزی بود. قبل از آن در سال ۲۰۰۲، فردی به نام گاستافسون نتایج مشابهی را ارائه داده بود. سرانجام گاش و همکاران (۲۰۰۶) توانستند نشان دهند که در استنباط بیزی هم می‌توان از این تبدیل استفاده کرد. نتایج حاصل از آن در قالب قضیه زیر منتشر شد.

قضیه ۲.۱.۴. فرض کنید برای هر $i = 1, \dots, n$ ، $\sum_j y_{ij} \geq 1$. اگر برای ϕ_i ‌ها پیشین‌های مستقل ناسره، $\pi(\phi_i) \propto 1$ ، و برای سایر عناصر پیشگو پیشین‌های سره مستقل از ϕ انتخاب شوند، آن‌گاه توزیع پسین حاصل از درست‌نمایی تبدیل چندجمله‌ای - پواسن با توزیع پسین حاصل از درست‌نمایی چندجمله‌ای معادل است.

برهان. فرض کنید $\alpha_i = \exp(\phi_i)$ ، $i = 1, \dots, n$. در نتیجه $\pi(\alpha_i) \propto \alpha_i^{-1}$. همچنین فرض کنید پیشین $\pi(\beta)$ برای زمانی باشد که پیشگو تابعی از β است. توزیع پسین کناری β حاصل از

$L_p(\phi, \beta)$ برابر با

$$\begin{aligned}\pi(\beta|\mathbf{y}) &\propto \pi(\beta) \prod_{i=1}^n \int_0^\infty \left(\alpha_i^{-1} \prod_j (\alpha_i g_{ij}(\beta))^{y_{ij}} \exp(-\alpha_i g_{ij}(\beta)) d\alpha_i \right) \\ &= \pi(\beta) \prod_{i=1}^n \left(\prod_j (g_{ij}(\beta))^{y_{ij}} \int_0^\infty \alpha_i^{\sum_j y_{ij}-1} \exp(-\alpha_i G_i(\beta)) d\alpha_i \right) \\ &\propto \pi(\beta) \prod_{i=1}^n \prod_j \left(\frac{g_{ij}(\beta)}{G_i(\beta)} \right)^{y_{ij}} \\ &= \pi(\beta) L_M(\beta)\end{aligned}$$

است، که در آن $g_{ij}(\beta)$ برحسب توابعی از β و $G_i(\beta) = \sum_j g_{ij}(\beta)$ است و به وضوح با پسین β حاصل از $L_M(\beta)$ یکسان است. $\square$

قضیه زیر سره بودن پسین بالا را در شرایط عمومی تر نشان می دهد.

قضیه ۳.۱.۴. فرض کنید $\sum_j y_{ij} \geq 1$ ، $i = 1, \dots, n$. اگر $\pi(\beta)$ سره باشد، آن گاه $\pi(\beta|\mathbf{y})$ نیز سره است.

برهان. فرض کنید Θ فضای پارامتر β باشد، بنابراین با استفاده از قضیه بالا داریم

$$\int_{\beta \in \Theta} \pi(\beta|\mathbf{y}) d\beta \propto \int_{\beta \in \Theta} \pi(\beta) L_M(\beta) d\beta < \infty$$

$\square$

توجه داشته باشید اگر از پیشین $\pi(\alpha_i) \propto \alpha_i^{\alpha_i-1}$ ($\alpha_i > 0$) در قضیه های بالا استفاده کنیم، آن گاه پسین حاصل دیگر معادل با $\pi(\beta) L_M(\beta)$ نخواهد بود، زیرا در این حالت در توان $G_i(\beta)$ عبارت $\sum_j y_{ij} + \alpha_i$ به جای $\sum_j y_{ij}$ قرار می گیرد. در این رساله با در نظر گرفتن $L_M \propto L_P$ با پیشگوی

$$\eta_{ij} = \alpha_j + \sum_{k=1}^l x_{kij} \beta_{kj} + f(u_{ij}).$$

میانگین مدل پواسن بعد از تبدیل MP به صورت $\gamma_{ij} = \exp(\alpha_j + \sum_{k=1}^l x_{kij} \beta_{kj} + f(u_{ij}) + \phi_i)$ است. همان گونه که قبلا هم بیان شد، برای مقایسه کاملتر و رسیدن به نتایج بهتر، مدل پیشنهادی را با مدل چندجمله ای کسری (لی، ۲۰۱۸) مقایسه می کنیم. برای این کار ابتدا مدل کسری را معرفی می کنیم.

۶.۱.۴ مدل چندجمله ای کسری

در این مدل متغیر پاسخ به عنوان کسری از وقایع اتفاق افتاده مورد نظر نسبت به همه حوادث رخ داده تعریف می شود. فرض کنید w_{ij} نسبت وقایع اتفاق افتاده از رده j ام برای آزمایش i

باشد، به گونه‌ای که $0 < w_{ij} < 1$ و $\sum_{j=1}^J w_{ij} = 1$. فرض کنید تابع w_{ij} کسری از (۴.۴) باشد به گونه‌ای که

$$E(w_{ij}|\eta_{ij}) = \frac{\exp(\eta_{ij})}{\sum_{s=1}^J \exp(\eta_{is})} = G_j(\eta_{ij}).$$

به روشنی معلوم است که $0 < G_j(\cdot) < 1$ و $\sum_{j=1}^J G_j(\cdot) = 1$. برای توسعه یک نسخه بیزی از مدل چندجمله‌ای کسری به شکلی که با مدل چندجمله‌ای پیشنهادی در بخش قبل قابل مقایسه باشد، برای این مدل توزیع‌های پیشین یکسانی را در نظر می‌گیریم. برای رفع مشکل شناسایی‌پذیری در این مدل نیز از قید رده گوشه و برای برازش مدل به داده‌ها به علت عدم امکان اجرای برازش با روش INLA از الگوریتم‌های MCMC استفاده می‌کنیم.

۲.۴ مطالعه شبیه‌سازی

در این بخش عملکرد مدل پیشنهادی را با یک مثال شبیه‌سازی با تولید اثر فضایی و احتمال‌های مربوط به رده‌ها مورد بررسی قرار می‌دهیم. به همین منظور نواحی $\{i = 1, \dots, n\}$ را بر روی یک شبکه منظم 10×10 تعریف کردیم تا به کمک آن یک مدل چندجمله‌ای فضایی را شبیه‌سازی کنیم. پیشگوی خطی

$$\eta_{ij} = \alpha_j + x_{ij}\beta_j + f(u_{ij}), \quad i = 1, \dots, n, \quad j = 1, \dots, J, \quad (7.4)$$

را در نظر گرفتیم، که در آن $n = 100$ تعداد نواحی و $J = 3$ تعداد رده‌های پاسخ هستند. اثر فضایی، $f(u_{ij})$ ، برای هر رده از مدل (۵.۴) با $\tau_u = 0.1$ پیروی می‌کند. برای تولید مشاهدات پاسخ از مدل لجیت چندجمله‌ای با پیشگوی خطی (۷.۴) استفاده کردیم که در آن β_j ها از توزیع نرمال استاندارد تولید و به گونه‌ای مقیاس‌بندی شدند تا جمع آن‌ها برابر صفر شود. همچنین $\alpha = (-1, -1, 2)$ تعیین شد. متغیر تبیینی x_j از توزیع یکنواخت استاندارد برای $j = 1, 2, 3$ شبیه‌سازی شد. بعد از استفاده از تبدیل چندجمله‌ای - پواسن، میانگین توزیع پواسن به صورت

$$\gamma_{ij} = \exp(\alpha_j + x_{ij}\beta_j + f(u_{ij}) + \phi_i),$$

است. از توزیع پیشین تعریف‌شده گاما در روش INLA برای τ_u و از توزیع‌های پیشین گوسی مبهم مستقل با میانگین صفر و واریانس بزرگ 10^5 برای $\{\alpha_j\}$ ، $\{\beta_j\}$ و $\{\phi_i\}$ استفاده کردیم. برای قید رده گوشه، $\alpha_1 = \beta_1 = 0$ و $f(u_{i1}) = 0$ تعیین شدند. همانند مثال مدل نرمال عاملی در فصل سوم، با داشتن برآوردهای سایر رده‌ها و با انجام محاسبات معمولی، می‌توان برآورد رده گوشه را نیز به دست آورد. فرض کنید β_2^* و β_3^* برآوردهای ضرایب رگرسیونی برای رده‌های دوم و سوم به ازای $\beta_1 = 0$ باشند. در این صورت

$$\beta_1 = \frac{-1}{3}\beta_2^* + \frac{-1}{3}\beta_3^*, \quad \beta_2 = \frac{2}{3}\beta_2^* + \frac{-1}{3}\beta_3^*, \quad \beta_3 = \frac{-1}{3}\beta_2^* + \frac{2}{3}\beta_3^*. \quad (8.4)$$

همچنین می‌توان روابط مشابهی را برای α_j و $f(u_{ij})$ یافت. برای قید مجموع مساوی با صفر قرار می‌دهیم

$$\sum_j f(u_{ij}) = 0, \quad \sum_j \beta_j = 0, \quad \sum_j \alpha_j = 0.$$

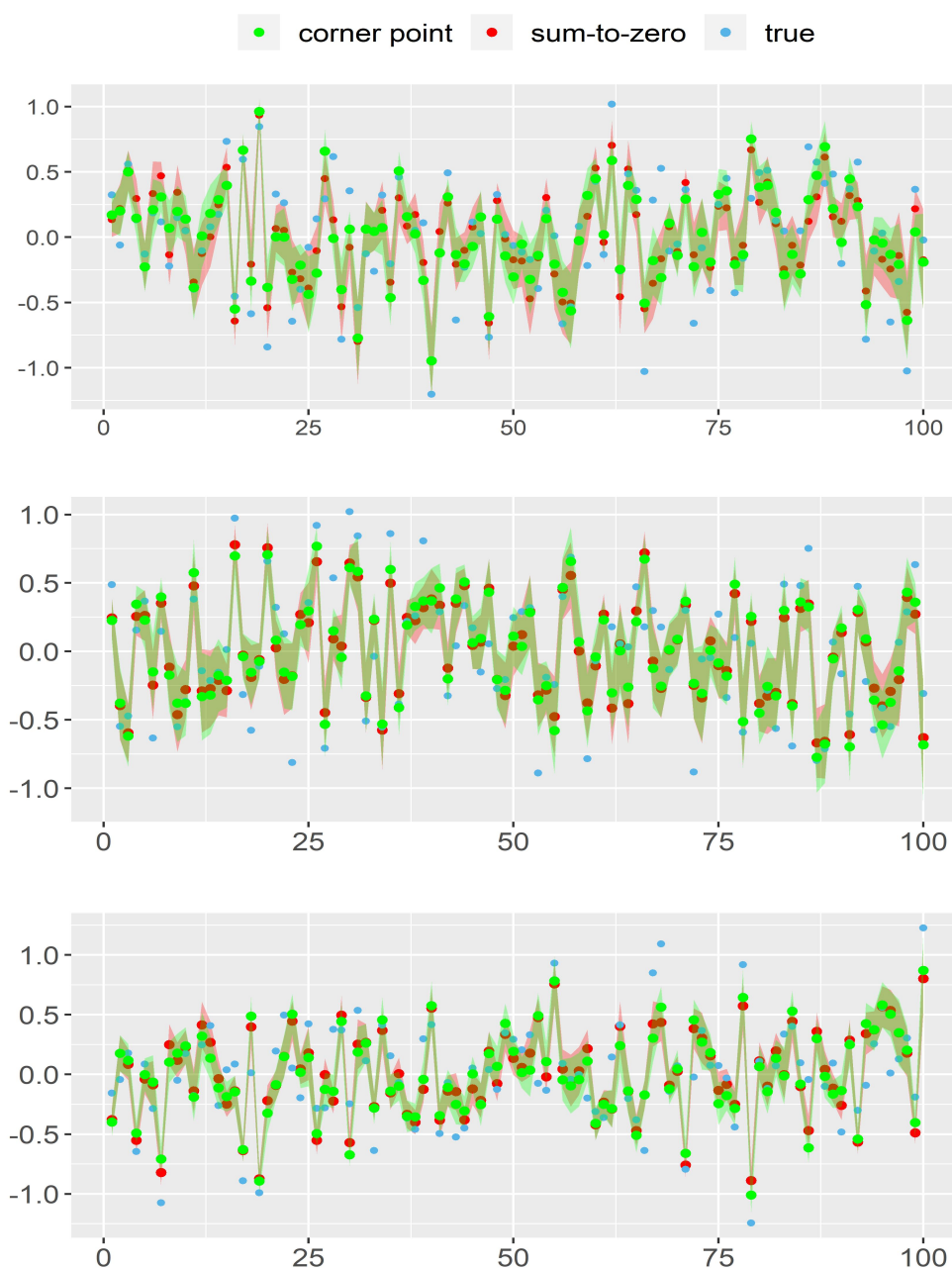
در هر دو این موارد، قید $\sum_i f(u_{ij}) = 0$ نیز به‌کارگرفته شد تا مدل CAR قابل شناسایی باشد. نتایج مربوط به برآورد اثر فضایی (شکل ۳.۴) در اکثر نواحی، تفاوت کمی را بین مقادیر برآوردشده و مقدار واقعی نشان می‌دهد. احتمال‌های برآوردشده در هر رده در شکل ۴.۴ برای هر دو قید شناسایی‌پذیری مدل چندجمله‌ای، مقادیر نزدیک به مقادیر شبیه‌سازی‌شده را نشان می‌دهد. چگالی‌های کناری پسین پارامتر τ_u برای هر دو نوع قید نیز در شکل ۵.۴ نمایش داده شده‌اند. برآورد τ_u برای قید مجموع مساوی با صفر تقریباً دقیق است، در حالی که برای قید رده گوشه بیش‌برآورد و حتی ناسازگار برآورد شده است.

جدول ۱.۴: برآورد میانگین (انحراف استاندارد) پسینی برای پارامترهای مدل با استفاده از دو قید شناسایی. از روابط (۸.۴) استفاده شده است تا برآورد پارامترهای رده گوشه بازیابی شوند.

پارامتر	مقدار واقعی	مجموع مساوی با صفر	رده گوشه
α_1	-۱	-۰/۹۰۴ (۰/۰۰۸)	-۱/۱۱۷ (۰/۰۰۹)
α_2	-۱	-۱/۰۹۷ (۰/۰۰۸)	-۰/۹۰۱ (۰/۰۰۹)
α_3	۲	۱/۹۷۲ (۰/۰۰۲)	۲/۰۲۰ (۰/۰۰۳)
β_1	-۰/۱۴	-۰/۱۳۱ (۰/۰۰۲)	-۰/۱۰۰ (۰/۰۰۳)
β_2	۱/۲۶	۱/۲۵۳ (۰/۰۰۵)	۱/۳۹۷ (۰/۰۰۶)
β_3	-۱/۱۱	-۱/۱۱۸ (۰/۰۰۶)	-۰/۹۸۶ (۰/۰۰۶)
τ_u	۰/۱۰	۰/۱۲۱ (۰/۰۱۱)	۲/۷۰۰ (۰/۲۸۰)

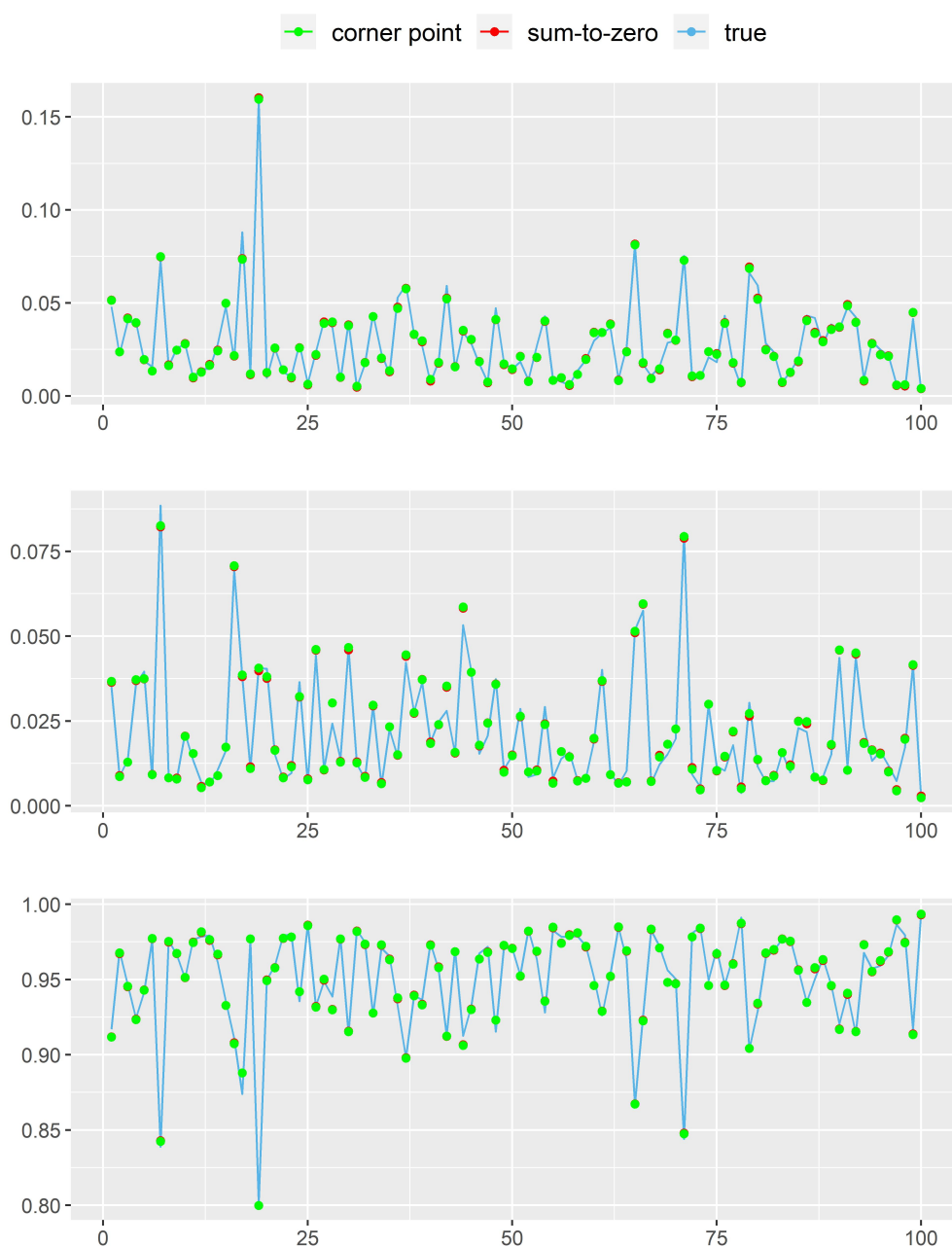
جدول ۱.۴ برآوردهای پسینی پارامترهای مدل شامل ضرایب رگرسیونی و پارامتر دقت اثر فضایی، τ_u ، را نشان می‌دهد. برآوردها برای قید مجموع مساوی با صفر تقریباً به مقادیر واقعی پارامترها نزدیک است. در مقابل، تفاوت بیشتری بین برآوردهای حاصل از به‌کارگیری قید رده گوشه با مقادیر واقعی به چشم می‌خورد.

در تفسیر قید مجموع مساوی با صفر همه پارامترها در تناسب با میانگین تفسیر می‌شوند که تقریباً دقیق برآورد شده‌اند. بنابراین انحراف استاندارد پسینی برای برآوردها کوچکتر خواهد بود. اثرات در رده گوشه با توجه به رده‌ای که برای این کار انتخاب می‌شود، برآورد می‌شوند. در واقع محققان یا آماردانان باید باورهای خود را نسبت به تفاوت‌های بین رده‌ها تعریف و تعیین کنند. در مقابل، روش مجموع مساوی با صفر هیچ رده خاصی را به‌طور ویژه در نظر نمی‌گیرد که خود موجب بهبود بخشیدن تفسیر می‌شود، زیرا مقدار متغیر پاسخ را در میانگین متغیرهای تبیینی بیان می‌دارد. در رده گوشه برای مثلاً رده z ، در تفسیر یک متغیر تبیینی با ثابت نگه داشتن سایر متغیرهای تبیینی، انتظار می‌رود با برآورد پارامتر مربوط عوض شود. در قید مجموع مساوی با صفر می‌توان برای هر رده توزیع پیشین تعریف کرد تا از مشکل تغییرپذیری در قید رده گوشه جلوگیری کرد.



شکل ۳.۴: میانگین پسین اثر فضایی در هر رده: نواحی سایه‌زده شده کران اعتبار بیزی ۹۵٪ را نشان می‌دهد. از بالا به پایین به ترتیب متعلق به رده‌های اول تا سوم است.

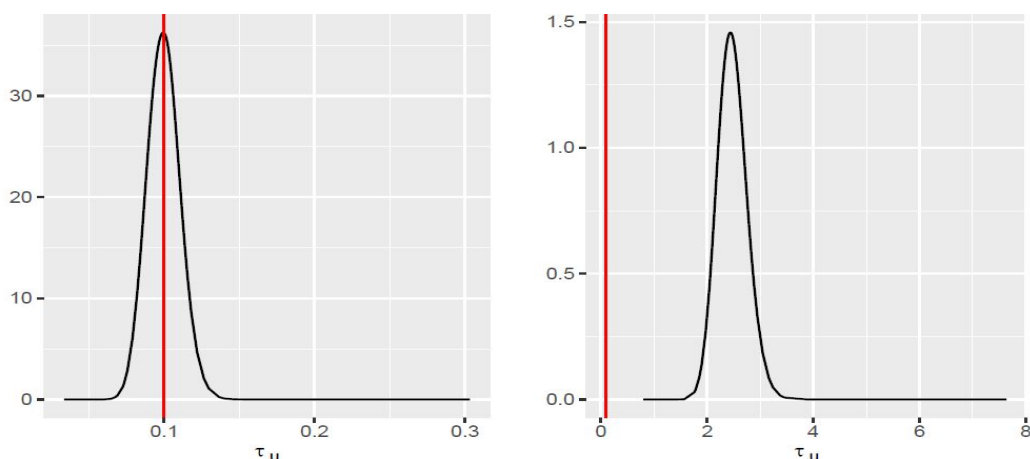
یافته‌ها در این قسمت با نتایج به‌دست آمده از شبیه‌سازی مدل نرمال عاملی در فصل سوم یکسان هستند.



شکل ۴.۴: برآورد احتمال هر رده: از بالا به پایین به ترتیب متعلق به رده‌های اول تا سوم است.

۳.۴ مثال کاربردی: شدت تصادفات جاده‌ای

امنیت در جاده‌ها یکی از بحث‌های کلیدی در بسیاری از کشورها است. تصادفات، علاوه بر خسارت‌های جانی و مالی خواه با مرگ همراه باشد یا نباشد، عواقب و هزینه‌های بعضاً جبران‌ناپذیری برای مردم و حکومت‌ها در پی دارند. ضمن این که ضربه‌های اجتماعی و روانی ناشی از تصادفات که حتی ممکن است سال‌ها بعد از آن ادامه یابد را نمی‌توان نادیده گرفت. بنابراین شناسایی عوامل ایجاد تصادف به‌ویژه در کشورهای پیشرفته توجه زیادی را به خود



شکل ۵.۴: چگالی‌های کناری پسین پارامتر دقت τ_u : قید مجموع مساوی با صفر (چپ)، و قید رده گوشه (راست). خط عمودی مقادیر واقعی را نشان می‌دهد.

جلب کرده است (سروری، ۲۰۲۰). این در حالی است که به بیان بانک جهانی تقریباً ۹۳ درصد از تصادفات در کشورهای در حال توسعه اتفاق می‌افتد (آمبو و همکاران، ۲۰۲۰). با این وجود، با یک فقدان بزرگ تحقیقاتی در این دسته از کشورها روبرو هستیم.

یک بخش کلیدی در تحلیل داده‌های مربوط به تصادف، کشف الگوهای مربوط به شدت جراحات و آسیب‌ها است (لیو و همکاران، ۲۰۲۰). پژوهشگران در بسیاری از کاربردهای مربوط به تحلیل داده‌های تصادف، مدل‌های گسسته را برای تحلیل شدت جراحات و آسیب‌ها به کار می‌گیرند (پنمتسا و همکاران، ۲۰۱۷). برای مثال، فارمر و همکاران (۱۹۹۷) اثرات مربوط به نوع ماشین و ویژگی‌های تصادف بر روی شدت جراحات را با کمک توزیع دوجمله‌ای مدل‌بندی کردند. دنل و همکاران (۱۹۹۶) یک پاسخ با چهار رده شدت جراحات را به کمک مدل لجیت ترتیبی و پروبیت ترتیبی تحلیل کردند. کالکمن و همکاران (۲۰۰۲) تحلیل تصادف‌های حاصل از برخورد دو ماشین به کمک مدل‌های احتمالی ترتیبی را بررسی کردند. همچنین، سای و همکاران (۲۰۰۹) از مدل بیزی پروبیت ترتیبی برای تحلیل داده‌های مربوط به جراحات ناشی از تصادف استفاده کردند. تعدادی از پژوهشگران از مدل لجیت چندجمله‌ای برای تحلیل این دسته از داده‌ها استفاده کردند که از جمله آن‌ها می‌توان به کارسون و همکاران (۲۰۰۱)، بریتمن و همکاران (۲۰۰۷)، شهید و همکاران (۲۰۱۳)، دانگ و همکاران (۲۰۱۳)، پنمتسا و همکاران (۲۰۱۷)، رضاپور و همکاران (۲۰۱۹) و وداگاما و همکاران (۲۰۲۰) اشاره کرد. لی و همکاران (۲۰۱۹) با در نظر گرفتن داده‌های پاسخ به صورت کسری از ماشین‌های تصادفی، مدل لجیت کسری را به کار گرفتند. از تحقیقات مربوط به این زمینه در سال‌های اخیر هم می‌توان به زنگ و همکاران (۲۰۱۹) اشاره کرد که در آن مدل لجیت ترتیبی تعمیم‌یافته برای مطالعه فراوانی تصادفات در آزادراه‌ها به کار گرفته شد.

مرور پژوهش‌های ذکر شده در بالا و مدل‌های به کار گرفته شده در آن‌ها به خوبی نشان می‌دهد که سطح‌بندی یا رده‌بندی میزان جراحات و آسیب‌های ناشی از تصادف همیشه مد نظر محققان بوده است، به‌طوری که تحقیقات فنگ و همکاران (۲۰۱۶) و بهنود و همکاران

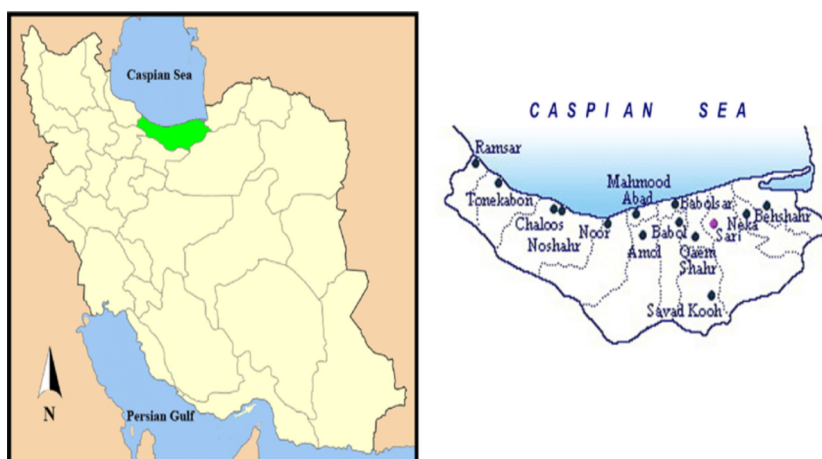
(۲۰۱۶) تاثیرات متفاوت ناشی از جراحات را خاطرنشان می‌کند. مدل لوجیت چندجمله‌ای خواه با رده‌های ترتیبی و خواه بدون ترتیب و همچنین مدل پروبیت ترتیبی برای مدل‌بندی سطوح مختلف شدت جراحات در داده‌های تصادف به دفعات فراوان مورد استفاده، نقد و بررسی قرار گرفته‌اند. ترتیب طبیعی شدت جراحات (از سطح بدون آسیب تا سطح آسیب جدی و مرگ) موجب شده است تا تحلیل‌گران متعددی از مدل لوجیت یا پروبیت ترتیبی به‌عنوان یک مدل احتمالی ترتیبی برای این دسته از داده‌ها استفاده کنند. یاماموتو و همکاران (۲۰۰۴)، ابدل (۲۰۰۳)، وئون و همکاران (۲۰۰۳)، کاتک و همکاران (۲۰۰۲)، کاتک و همکاران (۲۰۰۱)، کالک و همکاران (۲۰۰۲)، دائل و همکاران (۱۹۹۶)، الرو و همکاران (۲۰۱۲)، هائو و همکاران (۲۰۱۴) و زنگ و همکاران (۲۰۱۹)، تعدادی از محققان در این زمینه می‌باشند که از مدل‌های احتمالی ترتیبی استفاده کردند. با این وجود و برخلاف این محبوبیت، تعدادی محدودیت جدی بر این مدل‌ها حکمرانی می‌کند (ساولین، ۲۰۱۷). به‌طور مثال، تصادفات با شدت جراحات کمتر یا بدون آسیب‌دیدگی معمولاً گزارش نمی‌شوند که خود موجب اریبی در برآورد مدل‌های ترتیبی می‌شود (سالوم و همکاران، ۲۰۱۹). در مقابل، این مسئله در مدل‌های غیر ترتیبی ممکن است اثر ناچیزی را از خود به جای بگذارد (جانز و همکاران، ۲۰۱۳؛ سالوم و همکاران، ۲۰۱۹). علاوه بر این، برآوردهای رگرسیونی حاصل از مدل لجیت غیرترتیبی^{۱۳} سازگار هستند (مکفادن، ۱۹۸۱ و واشنگتن، ۲۰۰۳). بنابراین مدل چندجمله‌ای غیرترتیبی (اسمی) می‌تواند وجود برآوردهای سازگار از ضرایب رگرسیونی را برخلاف این عدم گزارش کامل تصادفات برای محققان تضمین کند. محدودیت دومی که محققان در استفاده از مدل‌های ترتیبی با آن روبرو هستند، تاثیری است که بر روی متغیرهای تبیینی دارد (سولین و منرینگ، ۲۰۰۷). به این معنی که تغییری که باعث افزایش احتمال در رده با بیشترین جراحت می‌شود، باعث کاهش احتمال مربوطه در رده با کمترین جراحت می‌شود (سولین و منرینگ، ۲۰۰۷؛ جیدایپالی و همکاران، ۲۰۱۱). بنابراین، استفاده از مدل‌های غیرترتیبی مثل مدل چندجمله‌ای برای تحلیل شدت جراحت در داده‌های تصادف پیشنهادی مناسب محسوب می‌شود اگرچه ذات این دسته از داده‌ها ترتیبی باشد (سالوم و همکاران، ۲۰۱۹). همچنین سالین و همکاران (۲۰۱۱)، مدل‌های مختلفی را مطالعه کرده و سرانجام یافتند که مفیدترین مدل برای مدل‌بندی گسسته داده‌های مربوط به شدت جراحات در تصادف، مدل لوجیت غیرترتیبی است.

برای اطمینان از این که داده‌ها قابل برازش به مدل چندجمله‌ای باشد، از آزمون استقلال جانشین‌های نامرتبط^{۱۴} (سالوم و همکاران، ۲۰۱۹) بهره گرفتیم. این ویژگی نشان می‌دهد نسبت احتمالات برای هر دو رده از شدت جراحات یکسان است اگرچه رده‌های دیگر در دسترس باشند یا خیر. در واقع این ویژگی یکی از مفروضات مهم توزیع چندجمله‌ای محسوب می‌شود. بنابراین یافتن این که ویژگی IIA برای داده‌های ما برقرار است یا نه امری ضروری

^{۱۳}Non-ordred multinomial

^{۱۴}Independence of irrelevant alternatives (IIA)

محسوب می‌شود که در این بخش مورد بررسی قرار خواهد گرفت. داده‌های تصادف به‌گونه‌ای اثر فضایی را از خود نشان می‌دهند که با کمک متغیرهای تبیینی قابل توصیف نیست. علاوه بر این، اثرات مربوط به متغیرهای تبیینی خود ممکن است به‌طور فضایی تغییر کنند. بنابراین، مدل‌های فضایی انتخاب طبیعی برای این دسته از داده‌ها محسوب می‌شوند (زیاکوپولوس و یانیس، ۲۰۲۰). در این رساله، مدل چندجمله‌ای بیزی سلسله‌مراتبی و از نوع فضایی را برای تحلیل داده‌های رده‌ای تصادف به کار می‌گیریم. با این هدف و با در نظر گرفتن همه جنبه‌های محل حادثه، شامل آسیب دیدن افراد و وسایل، سه رده از شدت آسیب و جراحات در تصادف را در نظر می‌گیریم. فهمیدن بهتر داده‌های تصادف به سیاستمداران این عرصه کمک می‌کند تا سیاست‌های مناسبی را برای کاهش روند تصادفات جاده‌ای انتخاب کنند. برای این کار داده‌های تصادف استان مازندران را مورد مطالعه قرار می‌دهیم (شکل ۶.۴ از پیردشتی و همکاران، ۲۰۱۵).



شکل ۶.۴: موقعیت جغرافیایی استان مازندران در نقشه ایران (قاب سمت چپ)، و نواحی مورد مطالعه (قاب سمت راست)

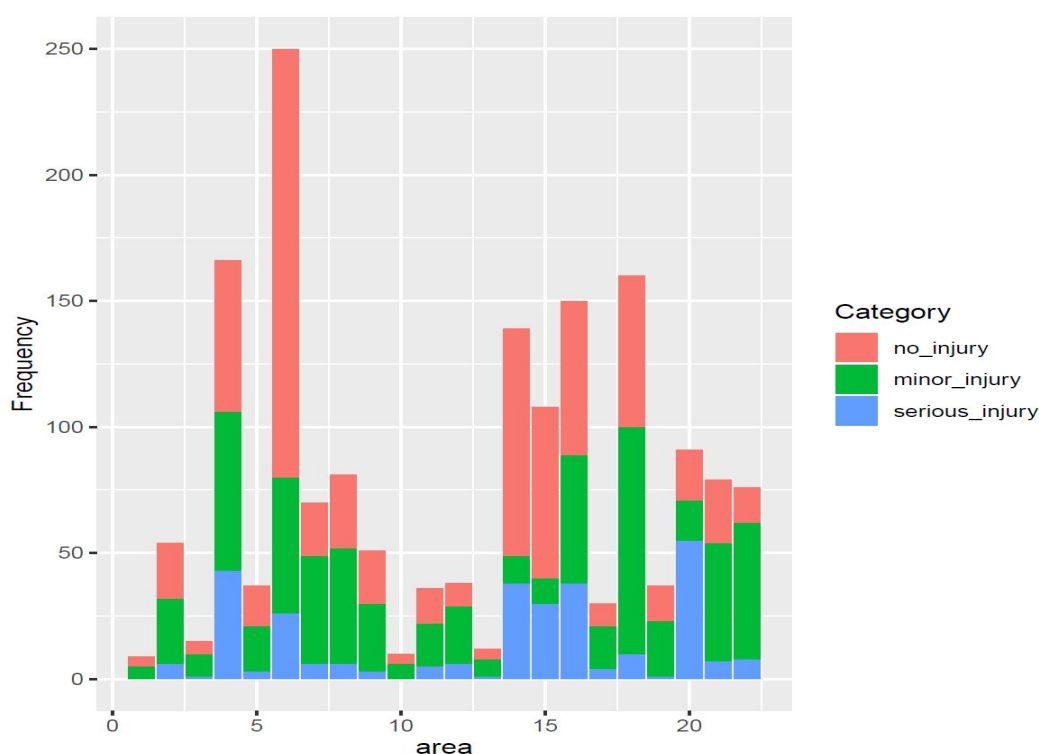
مازندران از نظر زمین‌شناسی و جغرافیایی، به دو بخش نواحی کوهی و جنگل‌های ساحلی تقسیم می‌شود. مساحت این استان ۲۳۸۴۲ کیلومتر مربع است. مازندران دارای ۲۲ مرکز فرمانداری است که به‌عنوان نواحی تصادف در نظر گرفته شده‌اند. مجموعه داده مورد نظر مربوط به تصادفات ثبت‌شده طی سال‌های ۲۰۱۶ تا ۲۰۱۸ در این مراکز است. در این مجموعه فقط تصادف‌های با اتومبیل در نظر گرفته شده‌اند. پس از چندی تلاش و پیگیری همه جانبه، به علت برخی مسایل امنیتی و نبودن اطلاعات ثبت‌شده کافی و جامع به تنها سه متغیر کمکی سرعت رانندگی، سن راننده (عامل تصادف) و نسبت تصادفات در روز نسبت به شب محدود شدیم. میزان جراحات افراد و خرابی ماشین‌ها و محیط پیرامون، به‌عنوان منبع تقسیم‌بندی شدت جراحات و رده‌بندی مورد استفاده قرار گرفتند. در نهایت سه دسته تصادفات به کمک مرکز پلیس به شرح زیر رده‌بندی شد:

۱. بدون آسیب (no injury): به تصادفی اشاره دارد که در آن افراد آسیب ندیده‌اند و خسارتی کمتر از ۳ میلیون تومان در پی داشته باشد.

۲. آسیب کم (minor injury): به تصادفی اشاره دارد که در آن افراد آسیب جدی ندیده یا با جراحات کمی مانند شکستگی‌های ساده روبرو شده‌اند. همچنین خسارات مالی ناشی از آن بین ۳ تا ۹ میلیون تومان باشد.

۳. آسیب جدی (serious injury): به تصادفی اشاره دارد که در آن افراد به‌طور جدی آسیب دیده یا با مرگ همراه شده‌اند. همچنین خسارت‌های مالی ناشی از آن بیشتر از ۹ میلیون تومان باشد.

از بین داده‌هایی که ما برای تحلیل و بررسی جدا کردیم، ۱۹٪ مربوط به آسیب‌دیدگی جدی، ۳۸٪ مربوط به آسیب کم و ۴۳٪ مربوط به تصادفات بدون آسیب است. فراوانی‌های مربوط به این رده‌ها در شکل ۷.۴ نمایش داده شده‌اند.



شکل ۷.۴: توزیع فراوانی مربوط به رده‌های جراحت برای همه ۲۲ ناحیه استان مازندران.

میانگین و انحراف استاندارد متغیرهای تبیینی انتخاب‌شده برای هر رده در جدول ۲.۴ گزارش شده‌اند. بنابراین با یک متغیر پاسخ سه رده‌ای و سه متغیر تبیینی روبرو هستیم.

جدول ۲.۴: میانگین (انحراف استاندارد) متغیرهای تبیینی در هر رده از پاسخ

متغیرهای کمکی	بدون آسیب	آسیب کم	آسیب دیدگی جدی
میانگین سرعت ماشین (km/h)	۵۶/۹۷ (۱۶/۶۸)	۶۷/۸۷ (۱۷/۱۹)	۸۹/۶۱ (۱۷/۶۵)
میانگین سن	۳۳/۰۶ (۷/۹۸)	۳۳/۹۹ (۹/۶۰)	۲۵/۷۰ (۱۰/۱۸)
نسبت تصادفات روز به شب	۰/۵۸ (۰/۳۰)	۰/۶۴ (۰/۲۵)	۰/۵۷ (۰/۲۳)

برای رسیدن به اطمینان از مفید بودن مدل پیشنهادی برای مجموعه داده مد نظر، برقرار

بودن ویژگی IIA را با کمک آزمون هاوسمن - مک‌فادن^{۱۵} (هاوسمن و مک‌فادن، ۱۹۸۴) بررسی کردیم. وقتی که این ویژگی برقرار باشد، برآوردهایی که با در نظر گرفتن تعدادی از رده‌ها به‌دست می‌آید، تفاوت معناداری با برآوردی که از همه رده‌ها با هم به‌دست آید، ندارد. بنابراین آزمون HM برآورد سازگار و کارای حاصل از مدل کامل، $\hat{\beta}^f$ را با برآورد سازگار اما ناکارا از مدل محدودشده، $\hat{\beta}^r$ را با استفاده از آماره

$$\chi_{HM}^2 = (\hat{\beta}^r - \hat{\beta}^f)' \{ \hat{V}(\hat{\beta}^f) - \hat{V}(\hat{\beta}^r) \}^{-1} (\hat{\beta}^r - \hat{\beta}^f),$$

بررسی می‌کند، که در آن $\hat{V}(\hat{\beta}^f)$ و $\hat{V}(\hat{\beta}^r)$ به ترتیب ماتریس کوواریانس‌های برآوردشده برای مدل کامل و مدل محدودشده هستند. در این‌جا χ_{HM}^2 فرض صفر این که پارامترهای مدل محدودشده با پارامترهای مدل کامل یکسان هستند را بررسی می‌کند که از توزیع کی‌دو با درجه آزادی برابر با متغیرهای برآوردشده در مدل محدودشده پیروی می‌کند. اگر فرض صفر رد نشود، یعنی وقتی که سطح معنی‌داری بیشتر از سطح آزمون (مثلاً ۵٪) باشد یا آماره آزمون χ_{HM}^2 منفی باشد، آن‌گاه ویژگی IIA برقرار است و مدل لجیت چندجمله‌ای مناسب است (زنگ، ۲۰۱۱). نتایج حاصل از آزمون HM در جدول ۳.۴ گزارش شده‌اند که به مناسب بودن مدل چندجمله‌ای برای مجموعه داده شدت تصادف‌های مازندران اشاره دارد.

جدول ۳.۴: نتایج آزمون HM برای بررسی ویژگی IIA در مجموعه داده تصادف‌های مازندران. علامت * سطح معنی‌داری برابر ۱ را نشان می‌دهد.

رده حذف‌شده	χ_{HM}^2	فرضیه صفر	ویژگی IIA
آسیب کم	-۵/۳۲	رد نشد	برقرار است
آسیب کم	۳/۰۶*	رد نشد	برقرار است
آسیب دیدگی جدی	-۲/۵۶	رد نشد	برقرار است

در مطالعات گذشته (برای مثال به رگرو و همکاران (۲۰۱۸) رجوع کنید)، نشان داده شده است سن اثر ناخطی با تصادفات دارد. علاوه بر این، رابطه همبستگی پیرسن بین سن و سرعت حدود ۲۱٪ برآورد شد که به نظر می‌رسد قابل اغماض باشد. به همین دلایل، به کمک فرآیند قدم زدن تصادفی، اثر متغیر تبیینی سن را با دو هدف به صورت ناپارامتری مدل‌بندی کردیم. هدف اول، دوری از هم‌خطی احتمالی بین سن و سرعت است. هدف دوم، لحاظ کردن رابطه ناخطی ممکن بین متغیر سن و متغیر پاسخ است. در آغاز، مدل را با همه متغیرهای تبیینی موجود در داده‌ها برازش دادیم. متغیر نسبت تصادفات روز به شب اثر معنی‌داری را از خود نشان نداد. بنابراین مدل را بدون حضور این متغیر تبیینی دوباره برازش دادیم. نتیجه دور از هم انتظار نیست، زیرا امروزه به‌دلیل پیشرفت در صنعت روشنایی جاده‌ای و به‌خصوص لامپ‌های مه‌شکن، لامپ‌های ایستاده با نور مناسب در خیابان‌ها به خوبی نصب شده‌اند.

^{۱۵}Hausman-McFadden (HM)

$$\eta_{ij} = \alpha_j + \text{speed}_{ij}\beta_j + f(\text{age}_{ij}) + f(u_{ij}), \quad i = 1, \dots, 22 \quad j = 1, 2, 3,$$

برای تحلیل نهایی به کار گرفته شد، که در آن $f(u_{ij})$ مدل فضایی ICAR و $f(\text{age}_{ij})$ اثر هموار سن را مدل‌بندی می‌کند. برای برآورد $f(\text{age})$ ، از فرآیند قدم زدن تصادفی مرتبه دوم استفاده کردیم. این فرآیند در عمل و کاربرد با یک مدل اسپلاین معادل است (رو و هلد، ۲۰۰۵). برای به کار گرفتن قید رده گوشه، رده بدون آسیب را به عنوان رده گوشه یا پایه انتخاب کردیم. توزیع‌های پیشین پارامترهای مدل را مشابه مثال شبیه‌سازی در نظر گرفتیم. برای تقریب توزیع پسین در مدل چندجمله‌ای کسری نیز از یک الگوریتم نمونه‌گیری متروپولیس-هستینگز (هستینگز، ۱۹۷۰) استفاده و با بسته mcmcpack (مارتین و همکاران، ۲۰۱۱) در نرم‌افزار R (تیم هسته مرکزی R، ۲۰۲۱) آن را اجرا کردیم. برازش مدل بیزی چندجمله‌ای کسری، حدود دو ساعت و نیم به طول انجامید، در حالی که اجرای مدل چندجمله‌ای با ساختار STAR با روش INLA کمتر از ۴ دقیقه طول کشید. همگرایی زنجیر MCMC تولیدشده برای مدل کسری نیز به کمک بسته نرم‌افزاری coda (کولس و کارلین، ۱۹۹۶) مورد بررسی قرار گرفت. نتایج برآورد مدل در جدول ۴.۴ گزارش شده‌اند. نتایج حاصل برای قید مجموع مساوی با صفر و مدل چندجمله‌ای کسری بیشتر قابل مقایسه هستند؛ به‌ویژه در برآورد پارامتر دقت τ_u . مشابه مطالعه شبیه‌سازی، قید رده گوشه عملکرد متفاوتی دارد که قابل اتکا نیست.

۱.۳.۴ بررسی مدل‌ها از نظر قدرت پیشگویی

توزیع پیشگو برای مشاهده جدید z برابر است با $\pi(z|\mathbf{y}) = \int_{\theta} p(z|\theta)\pi(\theta|\mathbf{y})d\theta$. توزیع پیشگو برای ارزیابی مدل از نظر قدرت پیشگویی مناسب است. برای ارزیابی توانایی پیشگویی مدل‌های ارائه‌شده، مقادیر مولفه پیشگوی شرطی^{۱۶} (پتیت، ۱۹۹۰) را برای مشاهدات می‌توان حساب کرد که بر اساس توزیع پیشگو محاسبه می‌شوند. برای مشاهده i ام مقدار CPO برابر

$$\text{CPO}_i = \pi(y_i|\mathbf{y}_{-i}),$$

است، که در آن $\mathbf{y}_{-i}$ بردار تمام مشاهدات بدون y_i است. در واقع، CPO مقدار احتمال پسین مشاهده y_i را بر اساس مدلی که بر روی همه مشاهدات بدون حضور مشاهده y_i برازش شده است، محاسبه می‌کند. مقادیر بزرگتر CPO نشان‌دهنده برازش بهتر مدل هستند. برای جمع‌بندی همه اطلاعات CPO ها در یک معیار عددی، می‌توان امتیاز لگاریتمی^{۱۷} را به کار برد. امتیاز لگاریتمی به صورت

$$\log.\text{score} = - \sum_i \log(\text{CPO}_i),$$

^{۱۶} Conditional predictive ordinate (CPO)

^{۱۷} Logarithmic score

محاسبه می‌شود. مقادیر کوچکتر این معیار کیفیت پیشگویی بهتری از مدل را نشان می‌دهند. امتیاز لگاریتمی مدل‌های برازش شده بر داده‌های تصادف‌های مازندران در جدول ۴.۴ گزارش شده‌اند. همان‌طور که انتظار داشتیم، مدل STAR با قید مجموع مساوی صفر از بقیه مدل‌ها برتر است و توانایی پیشگویی بالاتری دارد. همچنین در مقام مقایسه با مدل‌های دیگر، مقدار بزرگتر امتیاز لگاریتمی برای مدل چندجمله‌ای کسری، طبیعت ضعیف این مدل را در پیشگویی توصیف می‌کند.

۲.۳.۴ انتخاب مدل

برخی از اوقات ممکن است مدل‌های متعددی برای برازش و تحلیل داده‌ها در دسترس باشند به‌طوری که تشخیص واقعی مدل صحیح و مفید را با کمی پیچیدگی و چالش روبرو کند. در این‌گونه موارد به معیاری نیازمندیم که ما را در تعیین مدلی که مناسب داده‌های ما باشد، یاری رساند. به عبارتی، معیاری نیاز است که میزان سازگاری داده‌ها با مدل پیشنهادی را تعیین کند.

در مدل‌های بیزی، معمولاً افراد از معیار اطلاع انحراف^{۱۸} برای این کار استفاده می‌کنند. این معیار در ابتدا به صورت

$$D(\theta) = -2\log(p(y|\theta)),$$

تعریف می‌شود. اما در مدل‌های بیزی، معیار تعریف شده یک متغیر تصادفی محسوب می‌شود و از امید ریاضی آن تحت توزیع پسین به‌عنوان معیار برازش استفاده می‌کنند. با کمک ایده شمارش تعداد پارامترهای تأثیرگذار، می‌توان معیار

$$p_D = E(D(\theta|y)) - D(E(\theta|y)) = \bar{D} - D(\bar{\theta}),$$

را تعریف کرد. بنابراین

$$DIC = \bar{D} - p_D.$$

مقادیر کوچکتر DIC نشان‌دهنده مدل‌های بهتر هستند. برای داده‌های تصادف، در مقایسه با مدل‌های دیگر، مقدار کوچکتر DIC برای مدل با قید مجموع مساوی با صفر شاهد مدعی ما در برازش بهتر این مدل با این قید است.

همان‌گونه که از نتایج نشان داده شده در جدول ۴.۴ پیداست، متغیر تبیینی سرعت برای رده‌های آسیب کم و آسیب جدی، معنی‌داری است. بنابراین، سرعت زیاد احتمال تصادف با آسیب‌دیدگی را افزایش می‌دهد. این نتیجه با زنگ و همکاران (۲۰۱۶ و ۲۰۱۹) هماهنگ است. همچنین برآورد میانگین پسینی اثر ناپارامتری (ناخطی) سن با کران اعتبار بیزی ۹۵٪ در شکل ۸.۴، برای هر سه مدل و هر سه رده آسیب، نمایش داده شده است. نمودارها نشان از تغییر معنی‌دار اثر سن در رده‌های مختلف مربوط به شدت جراحات دارد. نتایج حاصل از

^{۱۸}Deviance information criterion (DIC)

هر سه مدل به‌طور تقریبی شبیه به هم هستند. شکل‌های مربوط به برآورد اثر سن، مناسب بودن مدل ناخطی ناپارامتری را برای برازش این اثر نشان می‌دهد.

به نظر می‌رسد افراد جوان‌تر باعث افزایش احتمال تصادف می‌شوند. اما افراد با سنین بالاتر تصادفات با خسارت کمتری را رقم می‌زنند که می‌توان علت را تجربه کافی در کنترل ماشین در شرایط سخت و ناگهانی به‌وجود آمده دانست. آن‌ها همچنین با جاده‌ها آشنایی بیشتری دارند. این نتایج با نتایج به‌دست آمده از محامد (۲۰۱۹) منطبق است. بنا بر پورتر (۲۰۱۱) در هر مایل مسافت پیموده‌شده، راننده‌های جوان با نرخ بالاتری تصادفات با مرگ و میر را موجب می‌شوند. بعد از دامنه سن جوانی، شیب مثبتی را برای بدون آسیب در سن‌های ۴۵ و ۶۰ سالگی می‌بینیم. به نظر می‌رسد این افراد تجربه لازم را برای کنترل ماشین در شرایط سخت دارند. آن‌ها، همچنین، قوانین مربوط به حداکثر سرعت در جاده‌ها را رعایت می‌کنند. بنابراین هنگامی که تصادف برای رانندگان بین ۴۵ تا ۶۰ اتفاق می‌افتد، احتمال این که خسارت ناشی از آن در رده بدون آسیب باشد، بیشتر است. از سویی دیگر، برای رانندگان با سن بیشتر از این، کاهش توانایی نگران‌کننده است زیرا ممکن است آن‌ها را به موقعیت‌های دشواری سوق دهد. اگرچه در داده‌های مربوط به استان مازندران رانندگان مسن‌تر از ۶۰ سالگی نداشتیم.

نتایج استنباط پسینی برای پارامتر دقت مدل ICAR، حضور اثر فضایی در مدل را تاکید می‌کند. عدم در نظر گرفتن آن موجب اریبی در برآوردها خواهد شد. میانگین و انحراف استاندارد مربوط به اثر فضایی به ترتیب در شکل‌های ۹.۴ و ۱۰.۴ نمایش داده شده‌اند. برآورد اثر فضایی برای هر سه مدل تقریباً مشابه است. همچنین، نتایج نرخ تصادف الگوی فضایی متفاوتی را در سرتاسر استان مازندران به رخ می‌کشد. بنابراین، یک نتیجه اساسی در این رابطه در نظر گرفتن اثر فضایی در تحلیل داده‌های تصادف مربوط به شدت جراحات است.

نقشه‌های مربوط به احتمال‌های برآوردشده در هر رده برای هر سه مدل، در شکل ۱۱.۴ نمایش داده شده است. نقشه‌ها نشان می‌دهند که احتمال تصادف در شهرهای توسعه‌یافته کم است. احتمال تصادف در شهرهایی که در مسیر اصلی تردهای بین استانی هستند، بیشتر است. بعضی از شهرها الگوهای مشابهی را از خود به نمایش گذاشته‌اند، مانند بابل، قائم‌شهر و ساری، که در آن‌ها احتمال داشتن تصادف با آسیب کم رو به افزایش است. تفاوت در کاربری زمین‌ها اثرات متفاوتی در میزان شدت تصادف دارند. در بعضی از شهرها افزایش فاصله از دریا، نزدیکترین جنگل و کوه احتمال تصادف را کاهش می‌دهد، به‌طوری که شهرهای با فاصله کمتر از دریا تصادفات بیشتری را تجربه می‌کنند. همچنین حضور گردشگران ناشی و نابلد در جاده‌های روستایی، احتمال تصادف با ماشین‌آلات کشاورزی را افزایش می‌دهد.

بنابراین در شهرهای با جاده روستایی بیشتر، احتمال تصادف نیز بیشتر است. عواملی مانند شرایط آب و هوایی ممکن است با قرار گرفتن در مدل، اثر معنی‌داری داشته باشند اما متأسفانه به دلیل در دسترس نبودن اطلاعات مربوط به آن‌ها در مدل مورد استفاده قرار نگرفتند. یافته‌ها ناهمگنی قابل توجهی را در داده‌های ما نشان می‌دهد و به نظر می‌رسد که مدل پیشنهادی آن را به خوبی بر عهده گرفته است. پس لازم است تا داده‌های مربوط به

تصادف برای تحلیل شدت جراحات را با مدل چندجمله‌ای برازش داده و اثر فضایی را در آن در نظر بگیریم. حضور توریست یا همان صنعت گردشگری در استان مازندران یکی از اساسی‌ترین صنایع در این استان محسوب می‌شود. با در نظر گرفتن اثر فضایی و توزیع احتمال در هر رده دریافتیم نواحی با احتمال حضور توریست بیشتر، تعداد تصادف بیشتری دارد. بنابراین در نظر گرفتن حضور توریست در تصمیمات مربوط به سیاست‌های پیشگیرانه تصادف امری ضروری است. تعطیلات تابستان احتمال حضور توریست را افزایش می‌دهد. از طرفی دیگر، عدم وجود آگاهی‌های لازم و نداشتن آشنایی کافی با جاده‌ها، به تصادفات با خسارت‌های مالی زیاد یا حتی مرگ منجر می‌شوند.

۴.۴ خلاصه فصل و نتیجه‌گیری

امروزه یادگیری ماشین^{۱۹}، هوش مصنوعی^{۲۰} و برنامه‌ریزی‌های صنعتی در علوم هم‌چون اقتصاد، علوم سیاسی و ژئوفیزیک نیاز به پیش‌بینی رده‌ها دارد که مدل لجیت چندجمله‌ای با اثر فضایی نقش مهمی را در این میان می‌تواند ایفا می‌کند. هم‌چنین داده‌های رده‌ای فضایی شبکه‌ای به علت پیشرفت در صنعت ثبت داده‌ها به‌طور گسترده‌ای در دسترس هستند که داده‌های تصادف که روی شدت جراحات رده‌بندی شده باشد، یکی از آن‌ها محسوب می‌شود. به همین دلیل، ما روش تحلیل داده‌های چندجمله‌ای فضایی با ساختار جمعی-رگرسیونی را انتخاب کردیم و از روش INLA به جای روش متداول نمونه‌گیری MCMC برای برازش مدل‌ها استفاده کردیم. برای آن‌که بتوان مدل را با INLA برازش داد، از تبدیل چندجمله‌ای-پواسن بهره‌مند شدیم. با استفاده از این تبدیل می‌توان به راحتی متغیرهایی را که به صورت یک تابع در مدل ظاهر می‌شوند، مانند اثرات ناخطی، اثر فضایی و اثر فضایی-زمانی افزون بر اثرات ثابت در پیشگوی جمعی-ساختاری مربوط به مدل چندجمله‌ای به کار گرفت.

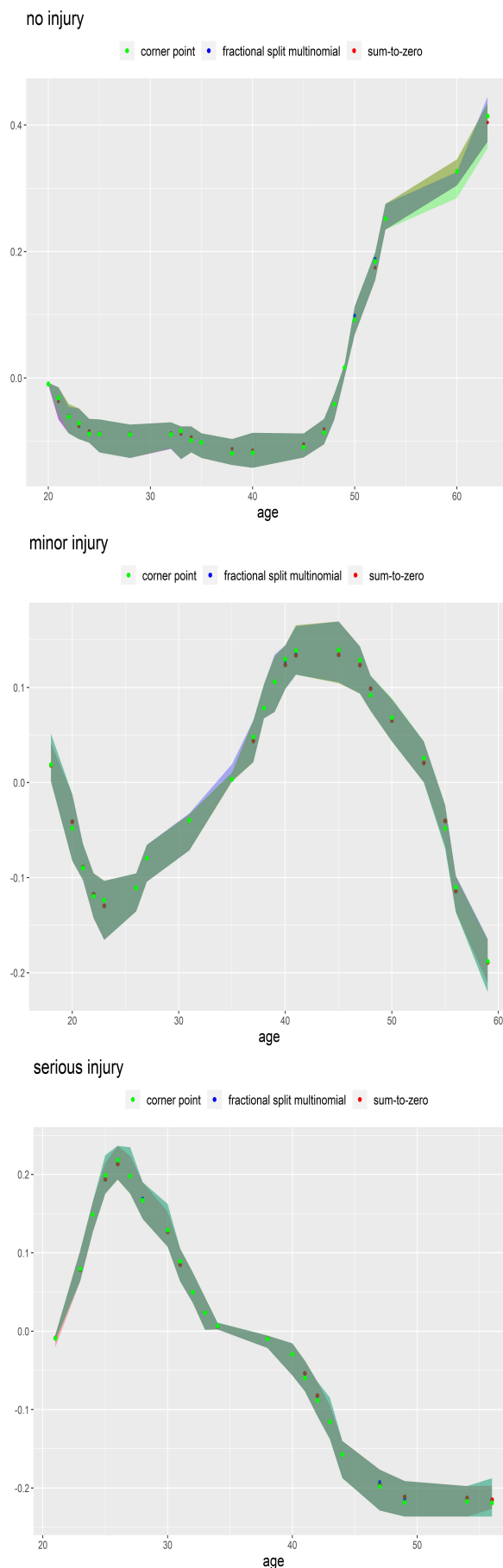
انجام استنباط بیزی در مدل چندجمله‌ای به دلیل وجود مشکل شناسایی‌پذیری، چالشی است. برای رفع این مشکل، قید مجموع مساوی با صفر را معرفی کرده و کارایی آن را در برآورد با قید متداول رده گوشه مقایسه کردیم. در همین حال مدل چندجمله‌ای کسری را نیز به علت مشابهت با مدل معرفی‌شده، به داده‌های واقعی برازش داده و نتایج حاصل از آن را مقایسه کردیم. خلاصه نتایج به شرح زیر است:

۱. با در نظر گرفتن دو معیار کارایی محاسباتی و توانایی پیشگویی، مدل چندجمله‌ای فضایی با ساختار جمعی-رگرسیونی و قید مجموع مساوی با صفر بر مدل چندجمله‌ای کسری برتری دارد.

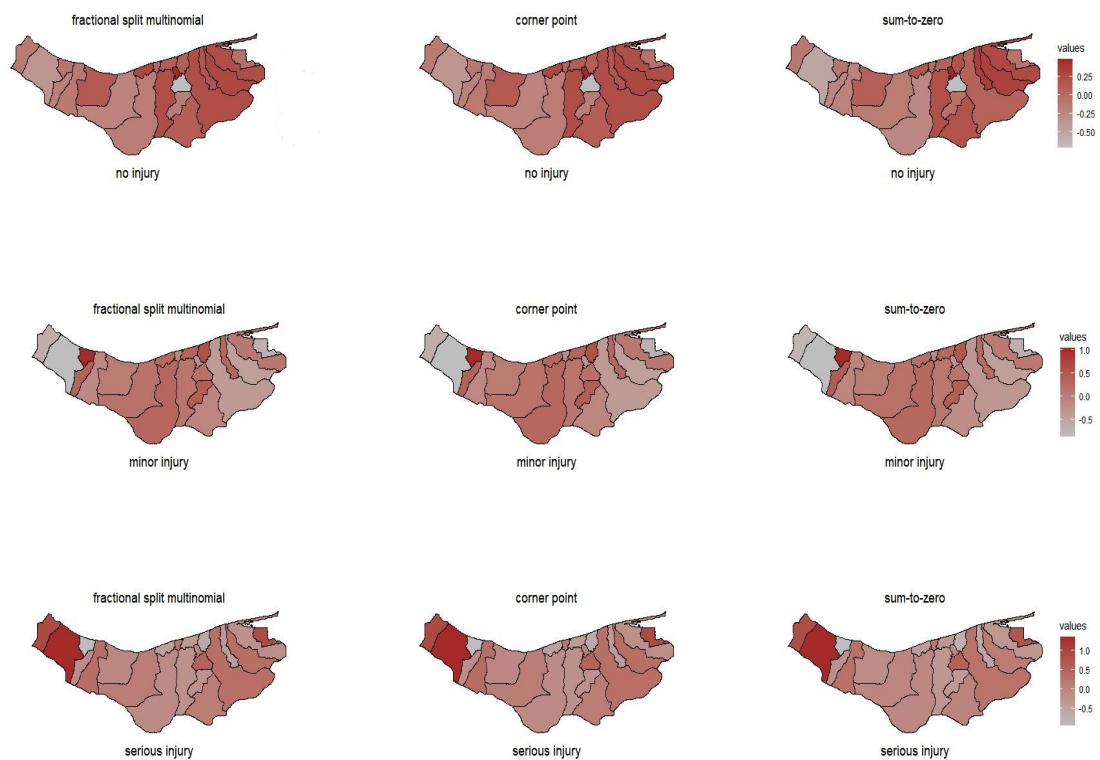
^{۱۹} machine learning

^{۲۰} artificial intelligence

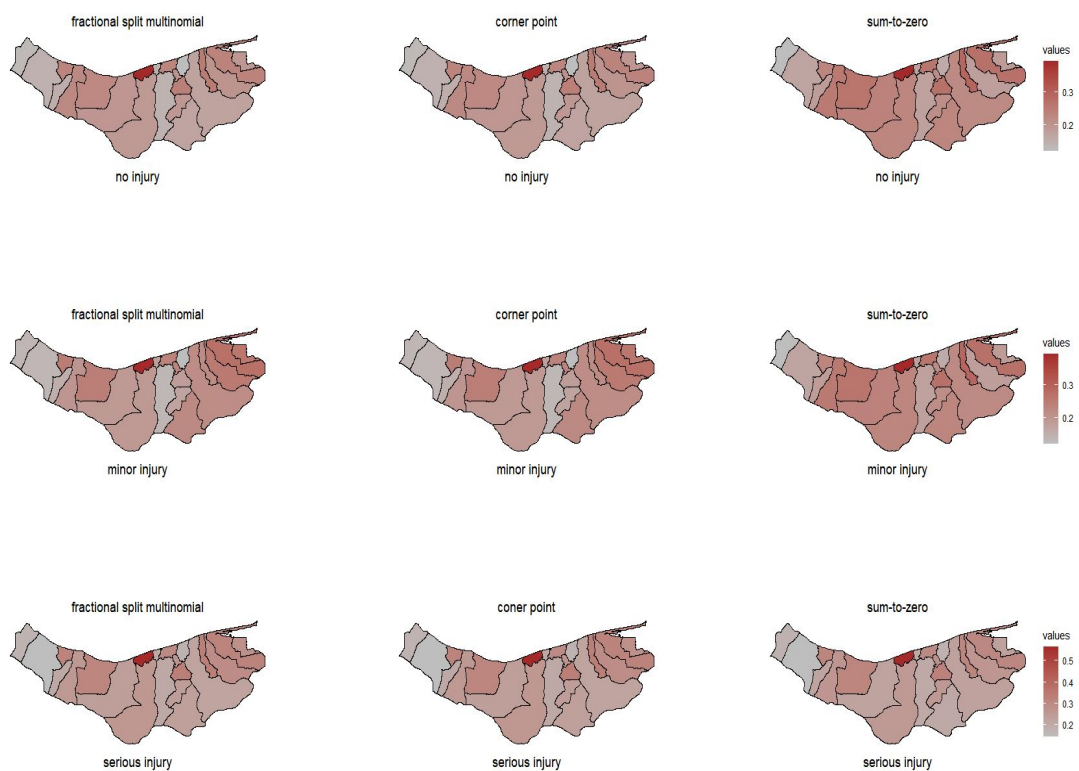
۲. مدل با قید مجموع مساوی با صفر برآوردهای کاراتری برای پارامترها به‌ویژه پارامتر دقت مدل بی‌سیگ فراهم آورده است.
۳. در حالت کلی تفاوت چشمگیری بین برآورد احتمال‌های رده‌ها دیده نمی‌شود و برآوردها تفاوت کمی با مقادیر واقعی دارند.
۴. اگر به آسانی تفسیر یک مدل باور داشته باشیم، باید مدل با قید مجموع مساوی با صفر را برگزینیم. زیرا این قید به بعضی از پارامترهای مدل نگاهی خاص ندارد؛ در مقابل، همه بر حسب تفاوت بین یک پارامتر و میانگین همان پارامتر در رده‌ها تفسیر می‌شود. بنابراین وقتی یک توزیع پیشین انتخاب می‌شود، می‌توان تفاوت بین پارامترها را حدس زد. در مقابل برای قید رده گوشه محقق باید برآورد را نسبت به رده انتخاب‌شده به‌عنوان رده گوشه تفسیر کند. از طرف دیگر، در قید مجموع مساوی با صفر توزیع پیشین را به همه عناصر موجود در رده‌ها نسبت می‌دهیم که این کار خود از مشکل تغییرپذیری توزیع پیشین جلوگیری می‌کند.



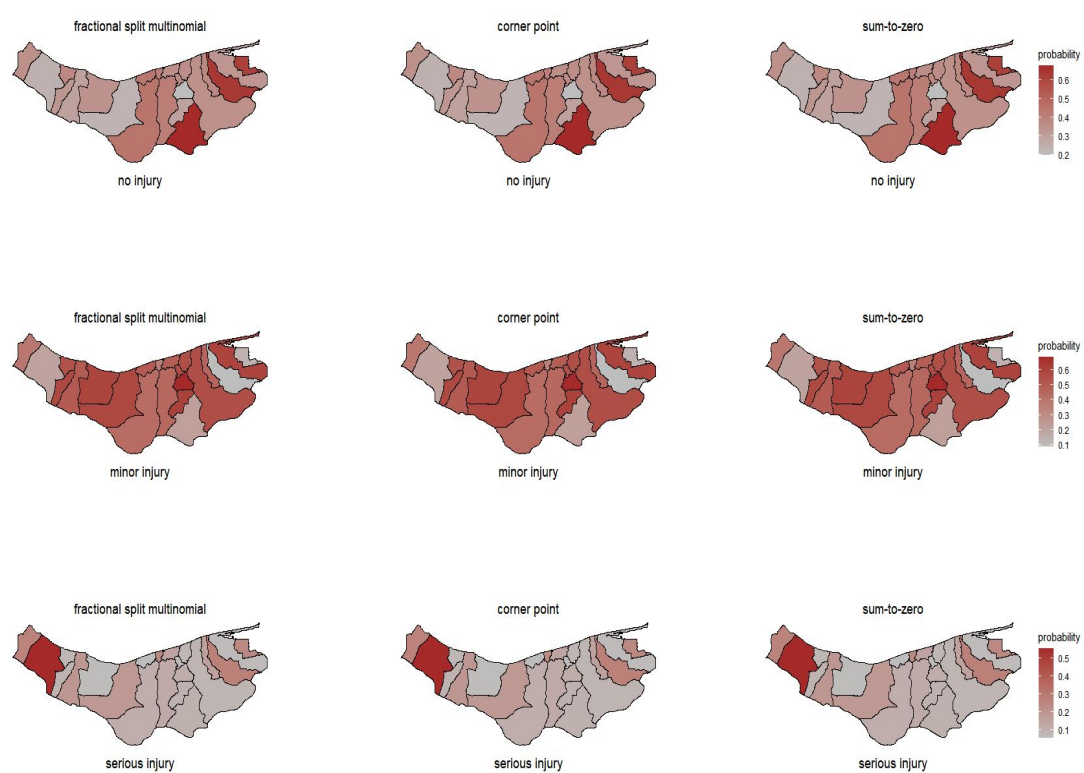
شکل ۸.۴: برآوردهای میانگین پسینی اثر سن با کران اعتبار بیزی ۹۵٪ در داده‌های تصادفات مازندران



شکل ۹.۴: برآوردهای میانگین پسینی اثر فضایی در هر رده و برای هر مدل: از بالا به پایین مربوط به رده‌های با شدت جراحت کم تا بالا.



شکل ۱۰.۴: برآورد انحراف استاندارد پسینی اثر فضایی در هر رده و برای هر مدل: از بالا به پایین مربوط به رده‌های با شدت جراحت کم تا بالا.



شکل ۱۱.۴: توزیع احتمال برآوردشده در هر رده برای هر مدل: از بالا به پایین مربوط به رده‌های با شدت جراحت کم تا بالا.

فصل ۵

مدل لوجیت تکی شده برای داده‌های چندجمله‌ای فضایی زمین‌آماري

در فصل قبل حالتی را بررسی کردیم که در آن متغیر پاسخ در هر رده دارای فراوانی بیشتر از یک است که برای داده‌های شبکه‌ای یک پذیره قابل قبول و استاندارد است. در این فصل داده‌های چندجمله‌ای را در نظر می‌گیریم که متغیر پاسخ فراوانی ندارد. در واقع، اثر فضایی را از نوع زمین‌آماري در نظر می‌گیریم، زیرا نبود فراوانی بیشتر از یک در این حالت، در اکثر کاربردها، معمول و پذیرفتنی است. برای رهایی از مشکل شناسایی پذیری و داشتن یک مدل با برآزش آسان و در عین حال تا حد قابل قبولی کارا، مدل لوجیت تکی شده^۱ (بگ و گری، ۱۹۸۴) را پیشنهاد می‌کنیم. به کمک شبیه‌سازی و یک مثال واقعی، کارایی برآوردهای حاصل از مدل پیشنهادی را با مدل چندجمله‌ای مقایسه می‌کنیم.

۱.۵ مدل لوجیت تکی شده

مدل (۱.۴) از فصل چهارم را در نظر بگیرید. فرض کنید متغیر Y_{ij} مقدار یک را بگیرد اگر پاسخ متعلق به رده z است و در غیر این صورت صفر. بر اساس مدل (۲.۴)، با این فرض که

^۱Individualized logit

متغیرهای تبیینی برای همه رده‌ها یکسان هستند، می‌توان رابطه

$$\log\left(\frac{\pi_{ij}}{\pi_{i1}}\right) = \mathbf{x}'_i \boldsymbol{\beta}_j, \quad i = 1, \dots, n, \quad j = 2, \dots, J \quad (1.5)$$

را نوشت، که در آن $\mathbf{x}_i = (x_{i1}, \dots, x_{il})'$ برای وارد کردن عرض از مبدأ، می‌توان قرار داد $x_1 = 1$. با پیروی از بگ و گری (۱۹۸۴) و اگوستی (۲۰۰۲)، اگر از یک مدل لوجیت چندجمله‌ای برای مقایسه رده j با رده اول استفاده کنیم، آن‌گاه رابطه

$$\log\left(\frac{\phi_{ij}}{\phi_{i1}}\right) = \mathbf{x}'_i \boldsymbol{\gamma}_j, \quad j = 2, \dots, J \quad (2.5)$$

برقرار است، که در آن

$$\phi_{ij} = P(Y_{ij} = 1 | \mathbf{x}_i, y_{i1} + y_{ij} = 1)$$

و

$$\phi_{i1} = P(Y_{ij} = 0 | \mathbf{x}_i, y_{i1} + y_{ij} = 1).$$

در این حالت تعداد $J-1$ مدل تکی شده تعریف می‌شود. در حقیقت در مدل لوژستیک تکی شده، مشکل تشخیص و شناسایی پارامترها به کمک یک مجموعه از مدل‌های لوجیت ساده به جای یک مدل چندجمله‌ای برطرف شده است. بنا بر بگ و گری (۱۹۸۴) می‌توان نشان داد دو مدل (۱.۵) و (۲.۵) از نظر پارامتربندی معادل هستند و برای $j = 2, \dots, J$ ، $\gamma_j = \beta_j$. به‌طور دقیق‌تر، با استفاده از قضیه بیز، می‌توان نوشت $\phi_{ij} = \frac{\pi_{ij}}{\pi_{i1} + \pi_{ij}}$. بنابراین

$$\frac{\pi_{ij}}{\pi_{i1}} = \frac{\phi_{ij}}{(1 - \phi_{ij})}.$$

در نتیجه، اگر با روش تکی شده $J-1$ رگرسیون لوژستیک برازش و پارامترها با استفاده از مدل (۲.۵) برآورد شوند، می‌توان از پارامترهای برآوردشده استفاده و با تعبیه آن‌ها در مدل (۱.۵) احتمال‌های رده‌ها را برآورد کرد.

به روشنی می‌توان درک کرد که روش تکی شده یا همان لوجیت تکی شده به صورت تحلیلی انعطاف‌پذیر است (رم و همکاران، ۱۹۹۵). در حالت کلی، روش تکی شده از کارایی بالایی برخوردار است (بگ و گری، ۱۹۸۴) و به‌ویژه برای مجموعه داده‌های بزرگ و تنک مفید است (رم و همکاران، ۱۹۹۵). این روش، همچنین برای توسعه یک چارچوب استنباطی بیزی با روش INLA نیز مناسب است.

۱.۱.۵ توزیع مجانبی برآورد درست‌نمایی ماکسیمم

تابع درست‌نمایی

$$L = \prod_{i=1}^n \prod_{j=1}^J \pi_{ij}^{y_{ij}},$$

را در نظر بگیرید. برآوردهای ماکسیمم درستنمایی حاصل از مدل چندجمله‌ای از $l(J-1)$ معادله

$$\sum_{i=1}^n x_{ki}(y_{ij} - \pi_{ij}) = 0, \quad k = 1, \dots, l, \quad j = 2, \dots, J$$

به دست می‌آیند، که در آن

$$\pi_{ij} = \frac{\exp(\eta_{ij})}{\sum_{s=1}^J \exp(\eta_{is})},$$

برآوردهای $\hat{\beta}_1, \dots, \hat{\beta}_J$ دارای توزیع نرمال مجانبی با میانگین $\beta_1, \dots, \beta_J$ و ماتریس کوواریانس $n^{-1}V^{-1}$ است، به گونه‌ای که $V = ((V_{jj'}))$ یک ماتریس $l(J-1) \times l(J-1)$ است که متشکل از زیرماتریس‌های $V_{jj'}$ با بعد $l \times l$ است و مولفه (k, t) آن‌ها عبارتند از

$$\lim_{n \rightarrow \infty} n^{-1} \sum_{i=1}^n x_{ki} x_{ti} \pi_{ij} (\delta(j, j') - \pi_{ij'})$$

که در آن دلتای کرونه کر است. در استفاده از مدل تکی شده، $J-1$ تابع درستنمایی شرطی ماکسیمم می‌شوند که هر درستنمایی به صورت

$$L_j = \prod_{C_j} \phi_{ij}^{y_{ij}} (1 - \phi_{ij})^{y_{i1}}, \quad j = 2, \dots, J,$$

است، که در آن $C_j = \{i : y_{i1} + y_{ij} = 1\}$. با قرار دادن $\ell_j = \log L_j$ ، می‌توان معادلات امتیاز متناظر را به صورت

$$\frac{\partial \ell_j}{\partial \beta_{kj}} = \sum_{i=1}^n x_{ik}(y_{ij} - (y_{ij} + y_{i1})\phi_{ij}) = 0, \quad k = 1, \dots, l, \quad j = 2, \dots, J,$$

نوشت. برآوردهای $\check{\beta}_j$ از حل معادلات امتیاز به دست می‌آیند. چون $\check{\beta}_j$ برآورد ماکسیمم درستنمایی است، بنا بر تعریف همگرایی در احتمال، می‌توان رابطه

$$\{\sqrt{m_j}(\check{\beta}_j - \beta_j) - m_j^{-1/2} H_j^{*-1} \frac{\partial \ell_j}{\partial \beta_j}\} \rightarrow 0$$

را نوشت، که در آن m_j اندازه (تعداد اعضای) مجموعه C_j را نشان می‌دهد و مولفه (k, t) ام ماتریس H_j^* عبارتست از

$$\lim_{m_j \rightarrow \infty} -m_j^{-1} \frac{\partial^2 \ell_j}{\partial \beta_{kj} \partial \beta_{tj}}.$$

در این جا می‌پذیریم که، نسبت برآمدها در هر رده به‌طور مجانبی، غیر قابل چشم‌پوشی است. چنین پذیره‌ای نامعمول نیست و به‌عنوان نمونه وقتی که متغیرهای تبیینی به‌طور یکنواخت کراندار باشند، برقرار است. با این پذیره، چون $m_j/n \rightarrow K_j > 0$ ، نتیجه

$$\{\sqrt{n}(\check{\beta}_j - \beta_j) - n^{-1/2} H_j^{-1} \frac{\partial \ell_j}{\partial \beta_j}\} \rightarrow 0,$$

به دست می‌آید، که در آن $H_j = K_j H_j^*$. چون این نتیجه برای همه j ها برقرار است، پس می‌توان نوشت

$$\{\sqrt{n}(\check{\beta} - \beta) - H^{-1} S_n\} \rightarrow 0 \quad (3.5)$$

که در آن

$$(\check{\beta} - \beta) = ((\check{\beta}_1 - \beta_1)', \dots, (\check{\beta}_J - \beta_J)'), \quad S_n = n^{-1/2} ((\partial \ell_1 / \partial \beta_1)', \dots, (\partial \ell_J / \partial \beta_J'))'$$

و H یک ماتریس قطری بلوکی با بلوک‌های $H_1, \dots, H_J$ است. با به کار گرفتن قضیه حد مرکزی روی S_n ، می‌توان همگرایی در توزیع زیر را نتیجه گرفت:

$$S_n \xrightarrow{D} N(0, W) \quad (4.5)$$

که در آن W یک ماتریس متشکل از زیرماتریس‌های $W_{jj'}$ با بعد $l \times l$ است که مولفه $W_{jj'}(k, t)$ آن‌ها عبارتست از

$$\lim_{n \rightarrow \infty} n^{-1} \text{cov} \left(\frac{\partial \ell_j}{\partial \beta_{kj}}, \frac{\partial \ell_{j'}}{\partial \beta_{tj'}} \right),$$

عناصر W را می‌توان با کمک رابطه

$$\text{cov} \left(\frac{\partial \ell_j}{\partial \beta_{kj}}, \frac{\partial \ell_{j'}}{\partial \beta_{tj'}} \right) = \sum_{i=1}^n x_{ki} x_{ti} \pi_i \phi_{ij} \phi_{ij'}^{-\delta(j, j')},$$

محاسبه کرد (بگ و گری، ۱۹۸۴). با استفاده از (۳.۵) و (۴.۵) می‌توان نتیجه گرفت که $\sqrt{n}(\check{\beta} - \beta)$ در توزیع به $N(0, H^{-1} W H^{-1})$ همگرا است. بگ و گری تاکید کردند که $H_j = W_{jj}$ و بنابراین بلوک‌های قطری $H^{-1} W H^{-1}$ برابر هستند با $H_j^{-1} = W_{jj}^{-1}$.

به‌طور خلاصه برآوردهای ماکسیمم درست‌نمایی حاصل از مدل چندجمله‌ای و مدل‌های لوجیت تکی شده، دارای توزیع مجانبی نرمال با بردار میانگین صفر و ماتریس‌های کوواریانس به ترتیب $V^{-1} n^{-1}$ و $H^{-1} W H^{-1}$ هستند. در حالت کلی مقایسه دو ماتریس کوواریانس راحت نیست، اما روشن است که عناصر خارج از قطر اصلی ماتریس $H^{-1} W H^{-1}$ که متناظر با کوواریانس‌های بین پارامترهای مدل‌های لوجیت جدا هستند، در حالت کلی ناصفر هستند.

۲.۱.۵ همگرایی برآوردگر بیزی مدل تکی شده

در این بخش، نشان می‌دهیم که برآوردگر بیزی پارامترهای رگرسیونی در مدل لوجیت‌های تکی شده به مقدار واقعی پارامترها همگرا است. برای این کار همگرایی در احتمال را با $\xrightarrow{P}$ نشان می‌دهیم.

قضیه ۱.۱.۵. در مدل لوجیت‌های تکی شده، چنان‌چه توزیع پیشین پارامترهای β روی فضای پارامتر Θ دارای چگالی مثبت و پیوسته باشد، آن‌گاه برآوردگر بیزی حاصل به مقدار واقعی پارامترها همگرا است.

برای اثبات این نتیجه به قضیه برنشتاین-وون میسس^۲ نیاز داریم که در زیر بیان شده است.

قضیه ۲.۱.۵. (باتاچاریا و همکاران، ۲۰۱۶). اگر توزیع پیشین π برای β دارای چگالی مثبت و پیوسته روی فضای پارامتر Θ باشد، آن‌گاه تحت π ، فاصله تغییرات کل^۳ بین توزیع پسین β و توزیع نرمال $N(\hat{\beta}_{ML}, -H^{-1})$ ، چنان‌چه $n \rightarrow \infty$ ، به صفر همگرا است، که در آن $\hat{\beta}_{ML}$ برآوردگر درست‌نمایی ماکسیمم پارامترها است.

برهان. با استفاده از نامساوی مثلثی و قضیه برنشتاین-وون میسس، برهان روشن است. فرض کنید $\hat{\beta}_{ML}$ و $\hat{\beta}_B$ به ترتیب برآوردگرهای ماکسیمم درست‌نمایی و بیزی پارامتر β باشند. بنابر نابرابری مثلثی می‌توان نوشت

$$|\hat{\beta}_B - \beta| \leq |\hat{\beta}_{ML} - \beta| + |\hat{\beta}_B - \hat{\beta}_{ML}|.$$

برای نشان دادن همگرایی مد نظر، باید نشان دهیم که سمت راست نابرابری در احتمال به صفر میل می‌کند. یکی از نتایج قضیه برنشتاین-وون میسس، همگرایی میانگین توزیع پسین در احتمال به برآوردگر ماکسیمم درست‌نمایی است. از طرفی دیگر، در بخش قبل دیدیم که برآوردگر ماکسیمم درست‌نمایی پارامترهای حاصل از مدل لوجیت‌های تکی شده برای β سازگار است. در نتیجه به سادگی نتیجه قضیه برهان می‌شود و می‌توان گفت $\hat{\beta}_B \xrightarrow{P} \beta$. □

۲.۵ مدل‌بندی وابستگی فضایی زمین‌آماري

تعریف داده‌های زمین‌آماري و میدان تصادفی گوسی را به خاطر بیاورید. علی‌رغم اهمیت و کاربردهای متنوع مدل‌بندی فضایی داده‌های زمین‌آماري با میدان‌های تصادفی گوسی، استفاده از آن‌ها برای موقعیت‌های کاربردی با حجم بالایی از داده‌ها، همیشه از نظر محاسباتی چالش‌برانگیز بوده است. چالش اصلی نیز مربوط به تجزیه ماتریس کوواریانس چگال میدان تصادفی است که هزینه محاسباتی آن از مرتبه $O(n^3)$ است. راهکارهای مختلفی برای حل این چالش محاسباتی پیشنهاد شده‌اند. برای دسترسی به یک مقاله مروری جدید در این حوزه و آشنایی با این راهکارها می‌توانید به هیتان و همکاران (۲۰۱۸) مراجعه کنید.

برای رهایی از این مشکل، لیندگرن و همکاران (۲۰۱۱) استفاده از روش معادلات دیفرانسیل جزئی تصادفی^۴ را پیشنهاد کردند. این روش بر اساس ابزارهای پیشرفته فرآیندهای تصادفی پایه‌ریزی شده است که از نظر محاسباتی کاراست و تاکنون استفاده‌های فراوانی از آن شده است.

^۲Bernstein–von Mises theorem

^۳Total variation distance

^۴Stochastic partial differential equation (SPDE)

میدان تصادفی گوسی کاملاً به وسیله میانگین و کوواریانس (ماتریس یا تابع) میدان شناسایی می‌شود. از نظر محاسباتی، کار کردن با یک میدان تصادفی گوسی، با بردارها و ماتریس‌های جبری روبرو است که برای آن می‌توان از بسته‌های استاندارد رایانه‌ای بهره گرفت. در مدل‌بندی آماری، افراد معمولاً از میدان تصادفی گوسی با ساختار کوواریانس پارامتربندی شده استفاده می‌کنند که اغلب زیرمجموعه‌ای از خانواده ساختار کوواریانس مترن است. روش SPDE از نمایش مارکوفی گوسی برای میدان تصادفی گوسی استفاده می‌کند که در آن بهره گرفتن از ماتریس دقت تنک، کارایی محاسباتی در روش‌های عددی را به‌ویژه در تجزیه ماتریس‌ها، به طرز قابل ملاحظه‌ای افزایش می‌دهد (رو و هلد، ۲۰۰۵). در این رساله از یک میدان تصادفی گوسی با میانگین صفر و تابع کوواریانس مترن بهره‌مند خواهیم شد. برای یادآوری، ضابطه تابع کوواریانس مترن در زیر نوشته شده است:

$$\text{Cov}(Y(s_i), Y(s_j)) = \frac{\sigma^2}{\Gamma(\lambda) 2^{\lambda-1}} (\kappa \|s_i - s_j\|)^{\lambda} K_{\lambda}(\kappa \|s_i - s_j\|)$$

که در آن، $\|s_i - s_j\|$ فاصله اقلیدسی بین موقعیت‌های $s_i, s_j \in D$ ، σ انحراف استاندارد کناری میدان مترن، $K_{\lambda}(\cdot)$ تابع بسل نوع دوم با مرتبه $\lambda > 0$ و $\Gamma(\cdot)$ تابع گاما است. پارامتر λ نیز درجه همواری فرآیند را اندازه‌گیری می‌کند و معمولاً ثابت نگه داشته می‌شود. همچنین κ پارامتر مقیاس است. به پارامترهای مدل مترن عنوان ابرپارامترهای مدل می‌دهیم و به میدان تصادفی گوسی پارامتربندی شده حاصل از آن، اثر تصادفی فضایی یا عنصر مدل فضایی زمین‌آماري نسبت می‌دهیم. از فواید استفاده از تابع مترن می‌توان به توانایی تفسیر فیزیکی پارامترهای آن اشاره کرد (باکا و همکاران، ۲۰۱۸).

راه‌های مختلفی برای انجام استنباط با ساختار کوواریانس میدان تصادفی گوسی وجود دارند. یک راه سنتی مرسوم ساختن ماتریس کوواریانس بر اساس مکان‌های مشاهدات مستقیماً از تابع کوواریانس و بعد ترکیب آن با ماتریس کوواریانس برای سایر عناصر مدل است تا یک ماتریس پرتبه برای مشاهدات یا مدل پنهان به وجود آورند. این روش برای مشاهدات تا اندازه صد موقعیت مکانی خوب کار می‌کند، اما برای داده‌های حجیم به اندازه صد هزار از نظر محاسباتی شدنی نیست، زیرا محاسبات با ماتریس کوواریانس Σ به شدت زمان‌گیر است. از طرف دیگر و فراتر از مسائل محاسباتی، ساختن تابع کوواریانس برای فضاهای هندسی به شکل کره (مثل سطح زمین) برای معرفی نامانایی در تابع کوواریانس و توسعه تابع کوواریانس فضایی در مدل‌های فضایی-زمانی جداناپذیر، چالش‌برانگیز است (باکا، ۲۰۱۸).

همان‌گونه که در فصل دوم مشاهده کردید، تمرکز روش INLA به‌جای ماتریس کوواریانس بر روی ماتریس دقت $Q = \Sigma^{-1}$ است، زیرا می‌توان نشان داد که مدل‌های مختلف با ساختار ماتریس کوواریانس پیچیده، ماتریس دقت تنک دارند (رو و هلد، ۲۰۰۵). ماتریس‌های دقت اکثراً صفر دارند و این صفرها در رایانه ذخیره نمی‌شوند. از طرفی در توزیع گوسی چندمتغیره، تنکی ساختار ماتریس دقت به استقلال شرطی بین متغیرهای تصادفی مربوط می‌شود (رو و هلد، ۲۰۰۵). فرض کنید بعد یک ماتریس دقت در یک مدل فضایی با ساختار فضایی تنک،

$n \times n$ باشد. استخراج نمونه، حساب کردن ثابت نرمال‌ساز و انجام محاسبات با ماتریس دقت به محاسباتی از مرتبه $O(n^{3/2})$ نیاز دارد و این در حالی است که هزینه انجام محاسباتی مشابه با ماتریس کوواریانس از مرتبه $O(n^3)$ است. از طرف دیگر، پژوهش‌های کاربردی که در آن‌ها صدها هزار مکان داشته باشیم با ماتریس دقت شدنی هستند و انجام دادن کارهایی پژوهشی فضایی با صدها مکان فقط به تعدادی ثانیه زمان نیاز دارد. روش SPDE یک نمایش از ماتریس کوواریانس مترن به‌عنوان پاسخی از معادله دیفرانسیل تصادفی

$$(\kappa - \Delta)^{\alpha/2} \tau u(s) = \Omega(s), \quad s \in D \subset \mathbb{R}^d \quad (5.5)$$

است، که در آن Δ عملگر لاپلاسی^۵ را نشان می‌دهد، τ واریانس را کنترل می‌کند، α درجه همواری را کنترل می‌کند، $\kappa > 0$ پارامتر مقیاس و $\Omega(s)$ یک فرآیند نوفه سفید گوسی فضایی با واریانس واحد است. ارتباط با پارامتر همواری ساختار مترن، یعنی λ ، و واریانس σ^2 از طریق $\alpha = \lambda + d/2$ و

$$\sigma^2 = \frac{\Gamma(\lambda)}{\Gamma(\alpha) (4\pi)^{d/2} \kappa^{\lambda/2} \tau^2}$$

است، که در آن d بعد میدان را نشان می‌دهد. پاسخ معادله را می‌توان با به‌کار گرفتن روش عناصر متناهی^۶ با توابع پایه و با نمایش مثلی با دامنه D به صورت

$$u(s) = \sum_{g=1}^G \phi_g(s) \tilde{u}_g \quad (6.5)$$

تقریب زد، که در آن G تعداد کل رأس‌های مثلی، $\{\phi_g\}$ مجموعه توابع پایه و $\tilde{u}_g$ وزن‌های با توزیع گوسی با میانگین صفر هستند.

بر اساس روش SPDE، می‌توان تقریب پیوسته‌ای از پاسخ معادله (5.5) که تقریباً ساختار کوواریانس مترن با دامنه کراندار محدودشده را دارد، ساخت. اگرچه پارامترهای معادله (5.5) در حالت استاندارد اما متفاوت از معادله (6.5) است، ولی رابطه یک به یک در میان آن‌ها وجود دارد. همچنین می‌توان از رابطه یک به یک بین پارامترهای تابع کوواریانس مترن و پارامترهای SPDE استفاده کرد تا با یک تقریب از نظر محاسباتی کارا به برآورد مدل پرداخت و در نهایت با پارامترهای معلوم تابع کوواریانس مترن نتایج حاصل را تفسیر کرد. برای مقادیر صحیح α ، پاسخ معادله SPDE مارکوفی است که در نتیجه ماتریس دقت تنک را برای تقریب‌های گسسته منعکس می‌کند. با در نظر گرفتن $\lambda = 1$ ، برای $d = 2$ ، خواهیم داشت $\alpha = 2$ و $\sigma^2 = \frac{1}{4\pi\kappa^2\tau^2}$. روشن است که دو پارامتر κ و τ اثری توأم بر واریانس کناری میدان تصادفی دارند.

سرعت رو به بالا اجرای مدل‌های فضایی در R-INLA، راحتی استفاده از مدل فضایی (به‌طور خاص داده‌های زمین‌آماري) همزمان با دیگر عناصر مدل و توانایی به‌کار گرفتن بخش وسیعی

^۵Laplacian operator

^۶Finite elements

از درستمایی‌های مشاهدات برای فرایندهای پنهان، موجب شده است تا INLA به ابزار خیلی مفیدی در استفاده از مدل‌های آماری کاربردی تبدیل شود. برازش مدل‌های خیلی پیچیده با روش INLA با رهیافت بیزی تنها ممکن است یک روز به طول انجامد. این امر موجب شده است تا پژوهشگران قادر باشند تا تعداد بیشتری از مدل‌ها را برازش دهند، به فهم بیشتری از داده‌ها برسند، درباره حساسیت پیشین‌ها بحث کنند، درباره حساسیت انتخاب درستمایی مشاهدات نظر دهند و به مقایسه پیشگویی‌ها بپردازند.

کاربردهای متنوعی از به کارگیری مدل‌های فضایی برای داده‌های زمین‌آماري در R-INLA وجود دارند که در زیر تنها به برخی از آن‌ها اشاره می‌کنیم. نور و همکاران (۲۰۱۴) در مجله لنست تحلیل فضایی-زمانی شدت انتقال مالاریا را مورد بررسی قرار دادند. گلدینگ و همکاران (۲۰۱۷) مرگ و میر زیر پنج سال و مرگ و میر نوزادان در چند کشور مختلف با یک مدل فضایی-زمانی جدپذیر در گروه‌های سنی مختلف را مدل‌بندی و تحلیل کردند. در مجله Science، جاسینو و همکاران (۲۰۱۷) اثرات تکه تکه شده در بیماری‌های عفونی پویا را مطالعه کردند. بات و همکاران (۲۰۱۵) اثر روش‌های مختلف در کنترل مالاریا را در آفریقا تحلیل کردند. همچنین، ظرفیت‌های INLA در مدل‌بندی فضایی یک ابزار کلیدی بود تا شادیک و همکاران (۲۰۱۸) برآورد سراسری قرار گرفتن در معرض محیط با ذرات بسیار ریز با قطر کمتر از $2/5$ میکرومتر را ارائه دهند که به $PM_{2/5}$ معروف است. نتایج مربوط به این پژوهش در سال ۲۰۱۶ در مطالعه جهانی بیماری (گاکیدو و همکاران، ۲۰۱۷) و ارزیابی سازمان بهداشت جهانی از خطر سلامتی ناشی از آلودگی محیط (سازمان جهانی سلامت، ۲۰۱۶) مورد استفاده قرار گرفت.

یک امتیاز بزرگ میدان تصادفی گوسی در روش SPDE با عناصر خطی، در مقایسه با مدل دقیق مترن، این است که انتگرال‌هایی از میدان تصادفی گوسی در یک زیرناحیه دلخواه از فضای مطالعه D را می‌تواند به عنوان ترکیب خطی از ضرایب نمایش دهد که به عنوان ردیف‌هایی از یک ماتریس بیان می‌شود. این به آن معنی است که می‌توان هم‌زمان، به صورت توأم، هم مشاهدات نقطه‌ای (زمین‌آماري) و هم مشاهدات شبکه‌ای را با هم برازش داد بدون این که در محاسبات مربوط به کوواریانس به مشکل برخورد (موراگا، ۲۰۱۷). از این ویژگی می‌توان، برای مثال، در مدل‌بندی هم‌زمان میزان بارش باران در ایستگاه هواشناسی و میزان آب روان از حوضه‌های آبریز رودخانه‌ها استفاده کرد. این روش همچنین، محاسبه عدم قطعیت را برای نقشه‌های کانتور مقدور می‌سازد، همان‌طوری که بولین و لیندگرن (۲۰۱۵) و بولین و لیندگرن (۲۰۱۷) نشان دادند. راحتی افزودن مشاهدات شبکه‌ای به این معنی است که در بعضی موارد (در صورت نیاز) می‌توان از این مدل به جای مدل بی‌سیگ بهره گرفت تا در مدل‌های فضایی با اندازه‌های در حال تغییر و مدل‌های فضایی-زمانی که مرزهای مربوط به نواحی در طول بازه زمانی مد نظر تغییر می‌کند، ساختار وابستگی واقعی‌تری را به نمایش گذاشت.

۱.۲.۵ پارامتربندی و پیشین‌ها

متأسفانه پارامترهای تابع مترن تفسیرپذیری مشکلی دارند. به همین علت لیندگرن و همکاران (۲۰۱۱) تعریف توزیع پیشین بر روی انحراف استاندارد کناری σ و دامنه تجربی $r = \sqrt{\lambda\nu/\kappa}$ را پیشنهاد و توصیه کردند. این توصیه در بسته نرم‌افزاری R-INLA نیز پیاده شده است. تفسیر این دو پارامتر راحت است. اگر بتوان چند نمونه از پارامترهای فضایی استخراج کرد، σ تغییرات میدان فضایی را نشان می‌دهد و r فواصل بین نواحی کوچک و بزرگ را متصل می‌کند. پیشین توأم برای r و σ به وسیله فلاگستمن و همکاران (۲۰۱۷) بر اساس رهیافت پیشین‌های با پیچیدگی جریمه شده^۷ (سیمپسون و همکاران، ۲۰۱۷) توسعه یافته است. توزیع پیشین پیشنهادی در عمل و کاربرد موفق بوده است (واکفید و همکاران، ۲۰۱۸). برای اطلاعات بیشتر در مورد این پیشین‌ها به سیمپسون و همکاران (۲۰۱۷) مراجعه کنید.

۳.۵ مقایسه روش‌های لجیت تکی شده و چندجمله‌ای به کمک شبیه‌سازی و مثال واقعی

۱.۳.۵ مطالعه شبیه‌سازی

در این بخش عملکرد مدل پیشنهادی را در مقایسه با مدل معمولی چندجمله‌ای در چارچوب تحلیل داده‌های زمین‌آماري به کمک شبیه‌سازی مطالعه و بررسی می‌کنیم. برای این کار، ۲۰۰ نقطه را در یک ناحیه $[0, 1] \times [0, 1]$ به صورت تصادفی تولید کردیم. پیشگوی جمعی به صورت

$$\mu_{ij} = \beta_{0j} + \beta_{1j}x_{1ij} + \beta_{2j}x_{2ij} + u(s_{ij}), \quad i = 1, \dots, 200, \quad j = 1, 2, 3$$

در نظر گرفته شد، که در آن $J = 3$ تعداد رده‌ها و $u(s_{ij})$ اثر فضایی است که تحقق‌های آن از یک میدان مترن با پارامترهای واقعی $\sigma^2 = 5$ ، $\nu = 1$ و $\kappa = 7$ تولید شدند. متغیرهای تبیینی x_1 و x_2 نیز از توزیع نرمال استاندارد تولید شدند. دقت کنید که تنظیمات شبیه‌سازی برای هر سه رده یکسان در نظر گرفته شد. رده اول را به عنوان رده گوشه انتخاب کردیم. بنابراین

$$\beta_{01} = \beta_{11} = \beta_{21} = 0, \quad u(s_{i1}) = 0.$$

با این قید مقادیر واقعی ضرایب رگرسیونی را به صورت $\beta_0 = (0, 1, 1)$ ، $\beta_1 = (0, -1, 1)$ و $\beta_2 = (0, 1, 1)$ در نظر گرفتیم. عملکرد روش تکی شده را با ارزیابی کارایی نسبی برآوردهای حاصل از آن نسبت به برآورد ماکسیمم درست‌نمایی محدود شده^۸ مدل چندجمله‌ای (وود،

^۷Penalized complexity (PC)

^۸Restricted maximum likelihood estimate

(۲۰۱۱) که در بسته نرم‌افزاری mgcv قابل محاسبه است، مورد بررسی قرار دادیم. در روش تکی شده، دو مدل لجیت شامل مقایسه رده دوم با اول و رده سوم با اول برآزش می‌شوند. توزیع‌های پیشین برای پارامترها و ابرپارامترهای مدل‌ها را همان پیشین‌های پیش فرض بسته R-INLA در نظر گرفتیم: یعنی پیشین‌های نرمال برای پارامترهای رگرسیونی با میانگین صفر و واریانس‌های بزرگ 10^3 و پیشین‌های PC برای ابرپارامترهای مدل مترن. برای ارزیابی کامل تر کارایی مدل، دو سناریو مختلف را مد نظر قرار دادیم. در سناریو اول حالتی را در نظر گرفتیم که در آن فراوانی مربوط به رده‌ها نامتعادل با نسبت ۰/۳۴ است. در سناریوی دوم، فراوانی مربوط به رده‌ها متعادل انتخاب شدند. جدول‌های ۱.۵ و ۲.۵ خلاصه‌ای از استنباط‌های پسینی مربوط به ضرایب رگرسیونی و پارامترهای تابع کوواریانس مترن را، برای دو سناریوی شبیه‌سازی، نمایش می‌دهند. علاوه بر این، میزان کارایی از دست‌رفته برآوردهای حاصل از مدل لجیت تکی شده را نسبت به مدل چندجمله‌ای حساب کردیم که در این جدول‌ها گزارش شده‌اند. برای محاسبه میزان کارایی از دست‌رفته از فرمول‌های

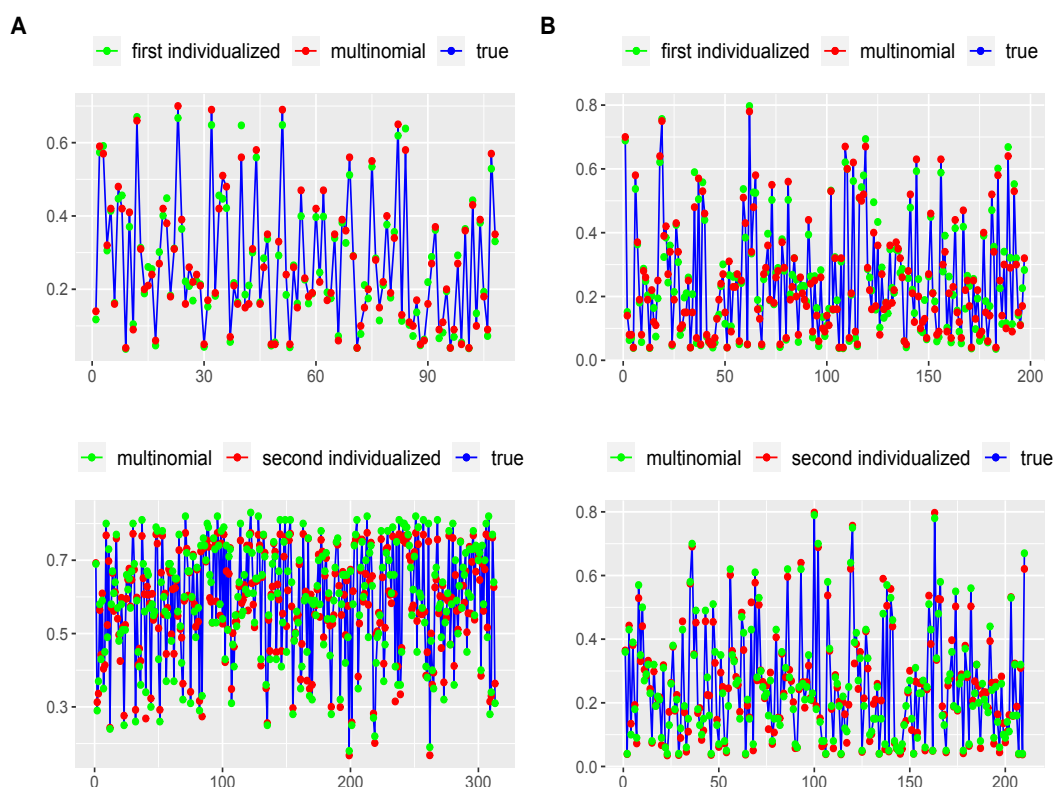
$$MSE(\beta) = \sum_{k=1}^K \frac{(\beta_{\ell j} - \hat{\beta}_{\ell j}^{(k)})^2}{K}, \quad \ell = 0, 1, 2,$$

$$MSE(\sigma) = \sum_{k=1}^K \frac{(\sigma - \hat{\sigma}^{(k)})^2}{K},$$

$$MSE(r) = \sum_{k=1}^K \frac{(r - \hat{r}^{(k)})^2}{K},$$

می‌توان استفاده کرد، که در آن β ، σ و r مقادیر واقعی هستند و $K = 100$ تعداد تکرارهای شبیه‌سازی است. همچنین $\hat{\beta}$ ، $\hat{\sigma}$ و $\hat{r}$ برآوردهای ماکسیمم درست‌نمایی محدودشده هستند که برای محاسبه MSE برآوردهای روش تکی شده باید آن‌ها را با برآوردهای بیزی (میانگین پسینی) جایگزین کرد.

برآوردهای حاصل از مدل بیزی لجیت‌های تکی شده و مدل کلاسیک چندجمله‌ای تقریباً به مقادیر واقعی نزدیک هستند. در حالت کلی، مقادیر مربوط به کارایی همه پارامترها بالا هست. اگرچه در بعضی موارد کارایی از دست‌رفته به حدود ۳۴٪ هم می‌رسد، اما دسترسی و آسانی برآزش مدل لجیت‌های تکی شده و دوری از مشکل شناسایی‌پذیری ذاتی مدل چندجمله‌ای، انگیزه کافی را برای نادیده گرفتن این میزان کارایی از دست‌رفته به خوبی ایجاد می‌کند. در مقایسه با دو سناریوی فراوانی متعادل و نامتعادل، روشن است که اگر فراوانی رده‌ها متعادل باشند، میزان از دست دادن کارایی برآوردهای مدل لجیت‌های تکی شده کمتر از حالت نامتعادل است. روند مشابهی در برآورد احتمال‌های رده‌ها دیده می‌شود. برآورد احتمال‌های رده‌ها برای دو سناریو در شکل ۱.۵ نمایش داده شده است که تقریباً به مقادیر واقعی نزدیک هستند.



شکل ۱.۵: برآورد احتمال رده‌های دوم و سوم در سناریوهای نامتعادل (چپ) و متعادل (راست). ردیف بالا مربوط به رده دوم و ردیف پایین مربوط به رده سوم است.

جدول ۱.۵: میانگین (انحراف استاندارد) پسینی ضرایب رگرسیونی در مدل‌های لجیت‌های تکی‌شده و چندجمله‌ای برای سناریوی نامتعادل.

پارامتر	مقدار واقعی	لوجیت تکی اول-دوم	لوجیت تکی اول-سوم	چندجمله‌ای		کارایی از دست‌رفته
				رده دوم	رده سوم	
β_{01}	۱	۰/۸۵ (۰/۰۸)		۰/۹۵ (۰/۰۹)		۰/۱۶
β_{02}	۱		۱/۱۹ (۰/۱۱)		۱/۰۱ (۰/۰۴)	۰/۲۴
β_{11}	-۱	-۰/۷۴ (۰/۰۵)		-۰/۹۱ (۰/۰۳)		۰/۳۴
β_{12}	۱		۰/۸۲ (۰/۱۱)		۰/۹۶ (۰/۰۸)	۰/۲۶
β_{21}	۱	۱/۲۰ (۰/۱۶)		۰/۹۵ (۰/۱۰)		۰/۲۸
β_{22}	۱		۱/۱۴ (۰/۱۲)		۰/۹۳ (۰/۰۵)	۰/۲۵
σ	۲/۲۳	۲/۲۱۹	۲/۵۴۰	۲/۱۲۲	۲/۱۴۴	۰/۱۵
r	۰/۴	۰/۳۵	۰/۴۳	۰/۴۳	۰/۴۱	۰/۲۲

۴.۵ مثال کاربردی: کیفیت آب‌های زیرزمینی

در این بخش، مدل‌بندی کیفیت منابع آب‌های زیرزمینی در استان گلستان مورد نظر است. استان گلستان در نقشه ایران در شکل ۲.۵ نشان داده شده است. مسئله کیفیت آب در سرتاسر

۹۰ مدل لوجیت تکی شده برای داده‌های چندجمله‌ای فضایی زمین‌آماري

جدول ۲.۵: میانگین (انحراف استاندارد) پسینی ضرایب رگرسیونی در مدل‌های لوجیت‌های تکی شده و چندجمله‌ای برای سناریوی متعادل.

پارامتر	مقدار واقعی	لوجیت تکی اول-دوم	لوجیت تکی اول-سوم	چندجمله‌ای		کارایی از دست‌رفته
				رده دوم	رده سوم	
β_{01}	۱	۰/۹۱ (۰/۱۰)		۱/۰۱ (۰/۰۴)		۰/۱۷
β_{02}	۱		۰/۹۸ (۰/۰۹)	۱/۰۵ (۰/۰۸)		۰/۱۰
β_{11}	-۱	-۰/۹۴ (۰/۰۷)		-۱ (۰/۰۳)		۰/۰۹
β_{12}	۱		۰/۹۲ (۰/۱۱)	۱/۰۱ (۰/۰۸)		۰/۲۳
β_{21}	۱	۱/۰۹ (۰/۱۱)		۰/۹۳ (۰/۱۰)		۰/۲۶
β_{22}	۱		۱/۰۰ (۰/۱۲)	۰/۹۵ (۰/۰۳)		۰/۱۵
σ	۲/۲۳	۲/۱۱	۲/۲۹	۲/۰۱	۲/۱۵	۰/۲۱
r	۰/۴	۰/۳۸	۰/۴۱	۰/۳۹	۰/۳۸	۰/۱۲

دنیا به‌ویژه در کشور عزیزمان ایران به‌دلیل کمبود آب و برخی از سایر مسائل جغرافیایی از اهمیت بالایی برخوردار است و همیشه نیاز به وجود آب تازه موجب شده است تا مردم گرایش فراوانی به سمت آب‌های زیرزمینی داشته باشند. هدف اولیه از تحلیل و مدل‌بندی این داده‌ها، پیدا کردن رابطه‌ای بین عناصر موجود در آب و کیفیت آن است تا بتوان اطلاعات مفیدی را برای مدیریت و برنامه‌ریزی آینده آب فراهم کرد.



شکل ۲.۵: استان گلستان: ناحیه هاشورخورد در نقشه ایران.

داده‌ها از ۱۵۸ موقعیت مکانی شامل چشمه‌ها و چاه‌های آب در بخش بزرگی از استان گلستان گردآوری شده‌اند. این موقعیت‌ها، در واقع، شامل ایستگاه‌های سنجش کیفیت آب هستند. استان گلستان به مساحت ۲۰۳۶۷ کیلومتر مربع در کنار دریای خزر واقع شده است. این استان هم از نظر کشاورزی و هم صنعتی از مناطق مهم کشور است.

برای اندازه‌گیری و سنجش کیفیت آب‌های زیرزمینی یک شاخص معمول و استاندارد، رسانایی الکتریکی^۹ است که با عبور جریان الکتریکی بین دو صفحه فلزی (الکتروود) اندازه‌گیری می‌شود. این متغیر بر حسب میکرومز^{۱۰} در هر سانتی‌متر اندازه‌گیری می‌شود و تحت تأثیر

^۹Electrical Conductivity (EC)

^{۱۰}Micromhos

مواد جامد مانند کلوراید، سولفات، منیزیم و سدیم است. در واقع هرچه مقدار این مواد در آب بیشتر باشند، مقدار EC بیشتر و در نتیجه کیفیت آب پایین‌تر است. بنا بر یک استاندارد پذیرفته‌شده جهانی، آب مناسب آشامیدن دارای مقدار EC کمتر از ۸۰۰ میکرومگ است (لیوز، ۲۰۰۳). به دلیل طبیعت آب‌های زیرزمینی و تغییر مقدار EC آن بر اساس اقلیم‌های متفاوت، وجود مواد جامد و املاح متفاوت در آن و تأثیر رفتار انسان‌ها بر آن، در نظر گرفتن تغییرات فضایی در تحلیل کیفیت آب امری ضروری است. اهمیت این ضرورت بیشتر هم درک می‌شود وقتی که بدانیم مدل‌بندی و تحلیل درست کیفیت آب بر اساس مدل‌های فضایی به نتایج زیر منتهی می‌شود:

- اطلاع‌رسانی و آگاهی‌بخشی به عموم مردم در استفاده از نقاط آب‌خیز برای آشامیدن یا مصارف دیگر.
 - تولید اطلاعات مفید برای مدیریت آب و برنامه‌ریزی‌های مربوط به آن
 - شناسایی مواد جامد با تأثیر زیاد بر میزان کیفیت آب
- با اهداف مطرح‌شده، برای رده‌بندی کیفیت آب از EC استفاده کردیم که در آن رده مشاهدات (متغیر پاسخ)، بر اساس لیوز (۲۰۰۳)، برای سه رده به‌صورت زیر تعریف شدند:

۱. آب قابل شرب: $EC < 800$

۲. کیفیت نامناسب برای شرب ولی مناسب برای کشاورزی: $800 < EC < 2500$

۳. کیفیت مناسب فقط برای اهداف صنعتی: $EC > 2500$

برای کنترل و مدیریت کیفیت آب باید در ابتدا عناصر موجود در آن را شناسایی کرد. بنابراین، باید اثرات مربوط به عناصری که در خاک اندازه‌گیری شده‌اند را مورد بررسی قرار دهیم که خلاصه‌ای از این عناصر در جدول ۳.۵ گزارش شده‌اند. مقادیر ثبت‌شده برای هر متغیر میانگین در سال است.

جدول ۳.۵: متغیرهای تبیینی موجود در مجموعه داده آب استان گلستان

متغیر تبیینی	نماد	واحد
پتاسیم	k	L ^۱ mg
بی‌کربنات	hco _۳	L ^۱ mg
کلوراید	cl	L ^۱ mg
سدیم	na	L ^۱ mg
سختی کل	th	L ^۱ mg
کل جامدات محلول شده	tds	L ^۱ mg

۱.۴.۵ مدل‌بندی و برآورد پارامتر

برای تحلیل داده‌ها یک مدل چندجمله‌ای فضایی داریم که پیشگوی جمعی در آن به صورت

$$\eta_{ij} = \beta_{0j} + \sum_{k=1}^6 \beta_{kj} x_{kij} + u(s_{ij}), \quad i = 1, \dots, 158, \quad j = 1, 2, 3$$

است، که در آن x_k ها، برای همه رده‌ها، همان متغیرهای تبیینی نوشته شده در جدول ۳.۵ هستند. رده سوم را به عنوان رده گوشه در نظر گرفتیم، زیرا این رده بیانگر کیفیت پایین آب است که نه برای آشامیدن مناسب است و نه برای کشاورزی. به همین دلیل مایل هستیم تا مقایسه‌ای که توسط مدل لجیت‌های تکی شده انجام می‌شود بین دو سطح کیفیت قابل شرب و قابل کشاورزی با رده‌ای باشد که فقط برای مصارف صنعتی مناسب است. توزیع‌های پیشین مدل لجیت‌های تکی شده را مشابه با مثال شبیه‌سازی انتخاب کردیم. با یک تحلیل مقدماتی در مورد انتخاب متغیرهای تبیینی موثر، تنها حضور دو متغیر سدیم و سختی کل در مدل نهایی تایید شد. نتایج برازش مدل کلاسیک چندجمله‌ای و مدل بیزی لجیت‌های تکی شده در جدول ۴.۵ گزارش شده‌اند. نتایج برازش در دو مدل تقریباً شبیه هم هستند؛ مقادیر انحراف استاندارد پسینی در مدل لجیت‌های تکی شده در اکثر موارد کوچکتر از انحراف استاندارد (مبتنی بر درست‌نمایی محدود شده) متناظر در مدل چندجمله‌ای است. در نتیجه فواصل اطمینان کلاسیک ۹۵٪ در مدل چندجمله‌ای کمی پهن‌تر از فواصل اعتبار بیزی مدل لجیت‌های تکی شده است. همان‌گونه که انتظار داشتیم ضرایب دو متغیر تبیینی سدیم و سختی کل در هر دو رده کیفیت قابل شرب و مناسب برای کشاورزی (نسبت به مصارف فقط صنعتی) منفی هستند که نشان‌دهنده تأثر معکوس این دو عامل بر کیفیت آب هستند؛ یعنی هر چه مقادیر این دو متغیر در آب بیشتر شوند، کیفیت آب کاهش پیدا خواهد کرد. شکل‌های ۳.۵ و ۴.۵، به ترتیب، نقشه‌های پهنه‌بندی اثرات فضایی و انحراف معیار متناظر را برای رده‌های اول و دوم در ناحیه مورد مطالعه نشان می‌دهند. نقشه احتمال وقوع هر رده نیز در شکل ۵.۵ نمایش داده شده است. نتایج حاصل از دو مدل قابل رقابت و شبیه هستند.

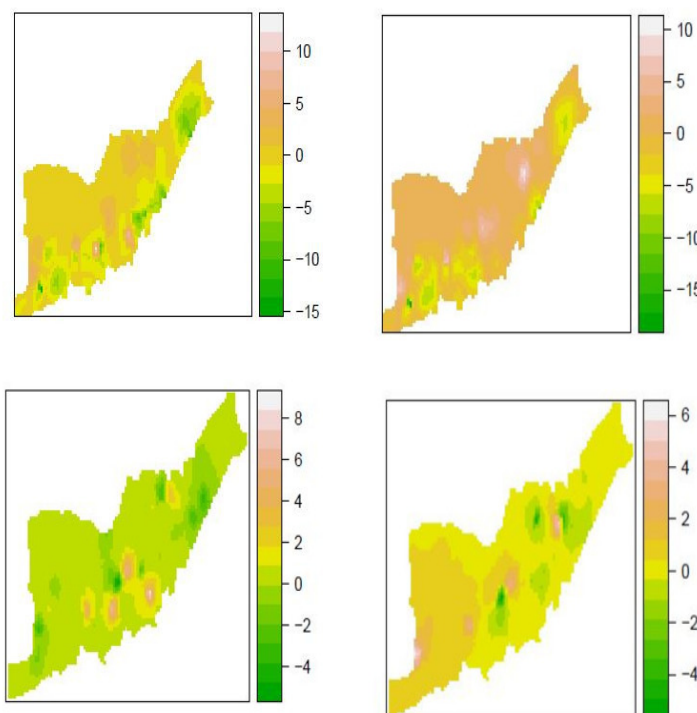
با در نظر گرفتن اثر فضایی در شکل ۳.۵ در می‌یابیم که در نواحی شمال شرقی و جنوبی اثر فضایی منفی است و کیفیت آب این مناطق برای شرب پایین است. در نواحی مرکزی و شمالی، آب با کیفیت قابل شرب را بیشتر می‌توان یافت زیرا این نواحی دارای اثر فضایی بزرگ و مثبت هستند. برآوردهای احتمال در شکل ۵.۵ نیز این نتایج را تایید می‌کنند.

۵.۵ خلاصه و نتیجه‌گیری

در این فصل تحلیل بیزی داده‌های چندجمله‌ای فضایی از نوع زمین‌آماري با مدل لجیت‌های تکی شده مد نظر قرار گرفت و عملکرد مدل پیشنهادی با مدل چندجمله‌ای کلاسیک مقایسه

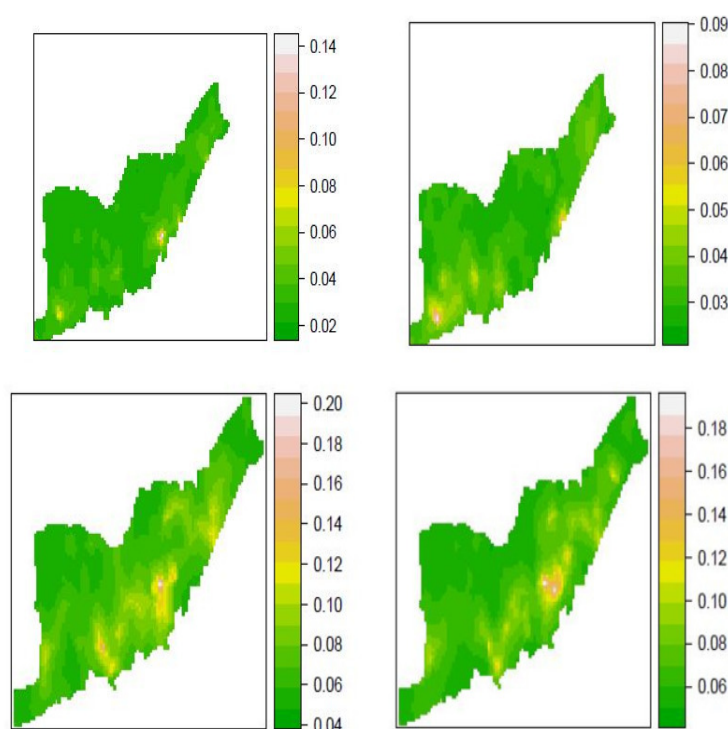
جدول ۴.۵: خلاصه‌ای از برآوردهای ضرایب رگرسیونی و پارامترهای کوواریانس مترن دو مدل برای مجموعه داده کیفیت آب استان گلستان.

مدل	رده اول			رده دوم		
	میانگین	انحراف استاندارد	CI ۹۵%	میانگین	انحراف استاندارد	CI ۹۵%
لوجیت تکی‌شده						
عرض از مبدأ	۳/۱۹	۰/۰۹	(۲/۱۳, ۴/۲۷)	۲/۶۴	۰/۰۶	(۰/۱۰, ۴/۲۲)
na	-۰/۰۳	۰/۰۸	(-۱/۰۱, -۰/۰۰۸)	-۰/۰۵۱	۰/۰۷	(-۰/۱۱, -۰/۰۰۲)
th	-۰/۰۶۲	۰/۱۶	(-۰/۸۱, -۰/۳۱)	-۰/۰۴۳	۰/۱۵	(-۰/۹۶, -۰/۰۱)
σ	۶/۲۳	۰/۲۱	(۲/۲۱, ۸/۵۹)	۷/۳۲	۰/۱۵	(۱/۲۳, ۱۳/۵۶)
r	۰/۷۱	۰/۱۵	(۰/۰۱, ۱/۳۹)	۱/۲۱	۰/۱۶	(۰/۶۳, ۲/۳۱)
چندجمله‌ای						
عرض از مبدأ	۴/۶۶	۰/۱۲	(۲/۵۵, ۶/۹۳)	۱/۸۲	۰/۱۱	(۰/۳۳, ۲/۵۱)
na	-۰/۰۱۱	۰/۱۱	(-۰/۱۵, -۰/۰۰۷)	-۰/۰۰۷	۰/۱۳	(-۰/۱۸, -۰/۰۰۳)
th	-۰/۰۶۹	۰/۱۹	(-۰/۹۱, -۰/۳۱)	-۰/۰۳۵	۰/۱۴	(-۰/۵۱, -۰/۱۲)
σ	۵/۲۴	۰/۱۲	(۱/۲۱, ۹/۴۱)	۶/۳۳	۰/۱۳	(۰/۰۱, ۱۲/۹۶)
r	۰/۶۸	۰/۳۱	(۰/۰۷, ۱/۳۱)	۰/۷۵	۰/۲۳	(۰/۱۶, ۱/۵۲)



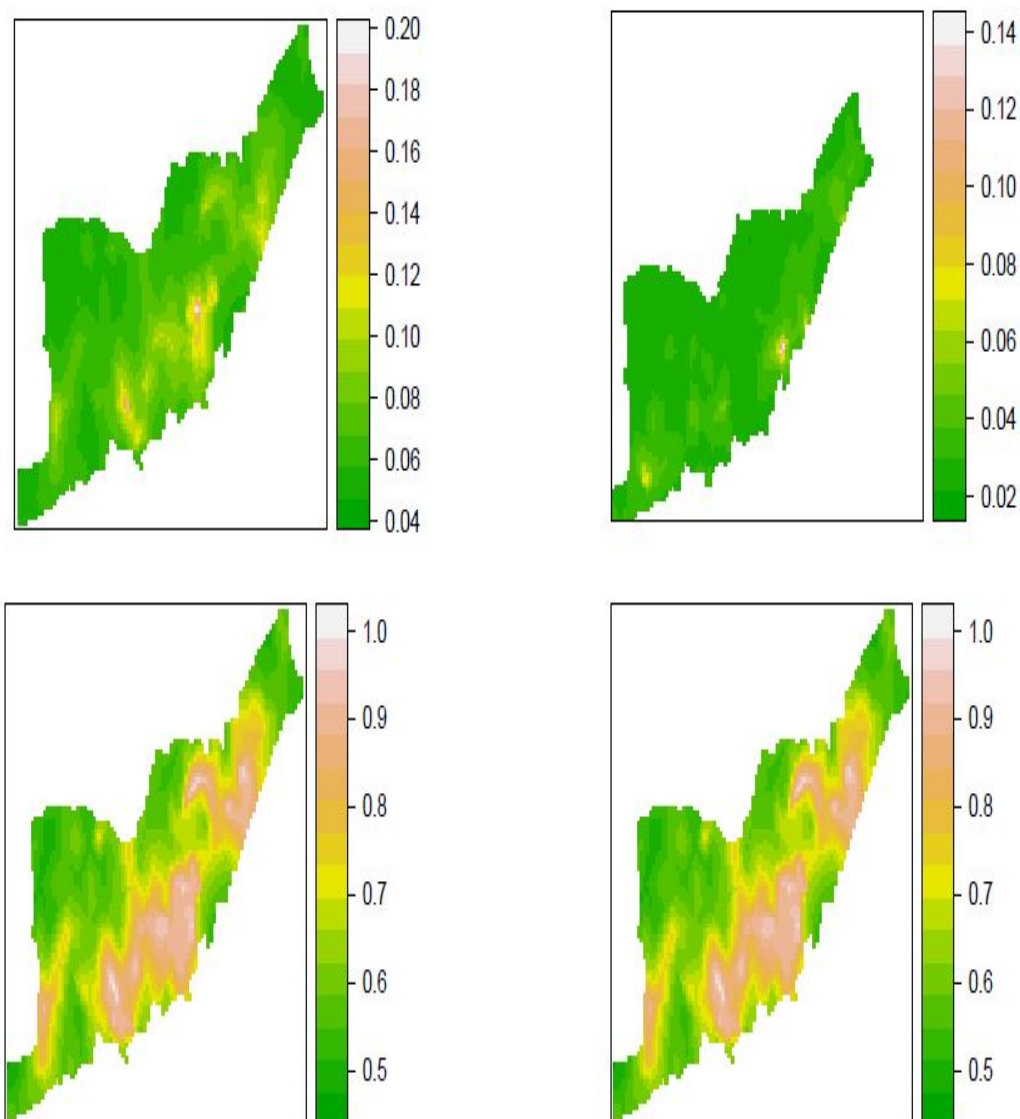
شکل ۳.۵: میانگین پسینی اثر فضایی برای رده اول (ردیف اول) و رده دوم (ردیف دوم) در دو مدل لوجیت‌های تکی‌شده (سمت چپ) و چندجمله‌ای (سمت راست).

شد. در واقع برای دوری از ساختار پیچیده مدل و داشتن یک مدل کارا با برازش آسان، مدل لوجیت‌های تکی‌شده را پیشنهاد کردیم. برای برازش مدل از روش INLA استفاده کردیم. با استفاده از این مدل‌بندی می‌توان حتی اثرات ناخطی را وارد مدل کرد و همچنان با سادگی برازش روبرو شد. از آنجایی که خیلی از بسته‌های نرم‌افزاری مربوط به روش بیزی قابلیت بهره‌مندی از مدل چندجمله‌ای با این ساختار را ندارند، یا زمان محاسباتی و حافظه مورد نیاز در آن‌ها حجیم است، یا قابلیت به‌کارگیری روش SPDE را ندارند به ناچار برای مقایسه



شکل ۴.۵: انحراف استاندارد پسینی اثر فضایی برای رده اول (ردیف اول) و رده دوم (ردیف دوم) در دو مدل لوجیت‌های تکی شده (سمت چپ) و چندجمله‌ای (سمت راست).

مدل پیشنهادی با مدل چندجمله‌ای از روش برآورد محدود شده در بسته نرم‌افزاری mgcv بهره گرفتیم. برای مقایسه جزئی‌تر نتایج مشاهده شده، کارایی برآوردها را نیز محاسبه کردیم. نتایج حاصل از دو مدل را نزدیک به یکدیگر و نزدیک به مقادیر واقعی یافتیم. میزان کارایی از دست‌رفته در مدل لوجیت‌های تکی شده با در نظر گرفتن سختی‌های پیش‌رو در برازش مدل چندجمله‌ای و دسترس نبودن بسته‌های نرم‌افزاری مناسب، قابل چشم‌پوشی است. از طرفی مشکل شناسایی‌پذیری ذاتی مدل چندجمله‌ای نیز وجود ندارد. در باور ما امروزه با پیشرفت در علومی همچون یادگیری ماشین و علم داده‌ها بسیاری از افرادی که لزوماً در رشته آمار نیستند، مثل متخصصان رشته‌های علوم اجتماعی، اقتصاد، و ژئوفیزیک، و دانش کدنویسی و ریاضی محدودی دارند، و حتی شاید توانایی لازم در به کارگیری مستقیم الگوریتم‌ها را ندارند، نیاز به تحلیل داده‌ها با این ساختار دارند که با به کارگیری مدل لجیت‌های تکی شده و افزایش کارایی آن به کمک روش INLA قادر خواهند بود تا بستری مناسب در این زمینه را فراهم آورند. در پایان نیز، برای ارزیابی بیشتر کارایی مدل، کاربست آن را برای تحلیل مجموعه داده واقعی مربوط به کیفیت آب در استان گلستان نشان دادیم.



شکل ۵.۵: برآوردهای احتمال برای رده اول (ردیف اول) و رده دوم (ردیف دوم) در دو مدل لوجیت‌های تکی‌شده (سمت چپ) و چندجمله‌ای (سمت راست).

فصل ۶

مدل‌بندی داده‌های چندجمله‌ای فضایی-زمانی

در این فصل سعی کردیم تا ساختار وابستگی فضایی-زمانی را در یک مدل چندجمله‌ای با ساختار جمعی-رگرسیونی تعبیه کنیم و داده‌های چندجمله‌ای فضایی-زمانی را با مدل پیشنهادی تحلیل کنیم. برای برآزش مدل از روش تقریبی INLA استفاده می‌کنیم. این نگاه برای مدل چندجمله‌ای، تا جایی که خبر داریم، تاکنون توسط سایرین استفاده نشده است. در این‌جا نیز برای رفع مشکل شناسایی‌پذیری مدل از قیدهای رده‌گوشه و مجموع برابر با صفر استفاده می‌کنیم. در ادامه، ابتدا نحوه وارد کردن انواع اثر متقابل فضایی-زمانی را تشریح می‌کنیم.

۱.۶ اثر فضایی-زمانی

امروزه پیشرفت ابزار و فنون ثبت داده‌ها و همچنین کارایی این ابزار در به‌روز رسانی و ذخیره کردن داده‌های حجیم با فاصله زمانی طولانی، موجب شده است تا افراد برای بررسی و تحلیل آماری پدیده‌هایی در علوم مختلف مانند پزشکی، بوم‌شناسی، زمین‌شناسی و علوم وابسته به آب و خاک که هم به زمان و هم به مکان وابسته‌اند، علاقه بیشتری از خود نشان دهند. در نتیجه، آن‌ها می‌توانند از انتشار دست‌آوردهای آماری برای برنامه‌ریزی‌های آینده کره زمین

بهره گیرند.

با توجه به جایگاه ویژه و اهمیت مدل چندجمله‌ای در تحلیل داده‌ها و در علوم مختلف، بر آن شدیم تا تحلیل کاراتر داده‌های چندجمله‌ای فضایی-زمانی را با در نظر گرفتن وابستگی فضایی-زمانی آن‌ها مد نظر قرار دهیم. برای اثر فضایی، مشابه فصل‌های قبل از مدل ICAR در یک ناحیه گسسته و مدل SPDE برای یک ناحیه پیوسته چگال می‌توان استفاده کرد. ساختار فضایی-زمانی در داده‌ها در فصل اول و بخش ۳.۶.۱ معرفی شد. برای یک مرور ساده می‌توان گفت، اثر فضایی را می‌توان به راحتی و با افزودن اثر زمانی به اثر فضایی-زمانی توسعه داد. یک فرآیند فضایی-زمانی به صورت

$$Y(s, t) \equiv \{y(s, t), (s, t) \in D \subset \mathbb{R}^2 \times \mathbb{R}^+\}$$

است، که در آن فرض می‌کنیم در n مکان فضایی (نقطه فضایی) و T نقطه زمانی مشاهده شده است. معمولاً افراد میدان تصادفی گوسی را برای این فرآیند در نظر می‌گیرند که دارای ماتریس کوواریانس $C((s, t), (s', t'))$ است. در این حالت، داده‌ها می‌توانند از نوع زمین‌آماري، شبکه‌ای یا نوع سوم یعنی الگوهای نقطه‌ای فضایی باشند. در این فصل، به‌طور خاص، داده‌های شبکه‌ای را در نظر می‌گیریم. در این حالت، مدل‌بندی دو ماتریس دقت مربوط به ساختار همسایگی مکان‌ها و ماتریس دقت مربوط به اثر زمانی مطرح هستند. مدل‌های مختلفی برای مدل‌بندی پدیده‌های فضایی-زمانی معرفی شده‌اند (ویکل و کرسی، ۲۰۱۱). راحت‌ترین راه، در نظر گرفتن اثر متقابل بین دو اثر فضایی و زمانی است. یکی از راه‌ها در نظر گرفتن اثر متقابل بین فضا و اتورگرسیو مرتبه اول در طول زمان است (مارتینز بنتیو و همکاران، ۲۰۰۸؛ ویوار و فرینو، ۲۰۰۹؛ راشورس و همکاران، ۲۰۱۴؛ رویز و همکاران، ۲۰۱۲؛ بلانجیاردو و همکاران، ۲۰۱۵). با این وجود، هیچ‌کدام از این مطالعات به فضا و زمان به‌عنوان اثر اصلی توجه نکرده‌اند. نور-هلد (۲۰۰۰)، مدلی را معرفی کرد که در آن اثرهای اصلی به قوت حضور دارند و برای مولفه فضایی-زمانی نیز چهار نوع اثر متقابل لحاظ می‌کند. در زیر جزئیات بیشتری از آن‌ها را تشریح می‌کنیم.

۱.۱.۶ اثرات متقابل ذاتی فضایی-زمانی

برای توضیح بیشتر درباره اثرات متقابل زمان و فضا، پیشگوی خطی را به صورت

$$\eta_{it} = \sum_{k=1}^l x_{kit} \beta_k + u_{it} + \nu_{it} + \delta_{it}, \quad i = 1, \dots, n, \quad t = 1, \dots, T, \quad (1.6)$$

در نظر می‌گیریم. همانند نور-هلد (۲۰۰۰) و بر اساس نیاز داده‌های در دسترس، مدل فضایی را مدل بی‌سیگ در نظر می‌گیریم. پیشگوی جمعی-ساختاری معرفی شده در (۱.۶) را در نظر بگیرید. مولفه u بیانگر اثر فضایی، ν نشان‌دهنده اثر زمانی و δ اثر فضایی-زمانی است. فرض می‌کنیم همه اثرات تصادفی دارای توزیع گوسی با میانگین صفر و ماتریس‌های دقت $\tau_u \mathbf{R}_u$

برای u ، $\tau_\nu R_\nu$ برای ν و $\tau_\delta R_\delta$ برای δ است، که در آن τ_δ ، τ_ν ، τ_u و R_δ ، R_ν ، R_u به ترتیب پارامترهای دقت و ماتریس‌های دقت متناظر با اثرات فضایی، زمانی و فضایی-زمانی هستند. در اینجا، R_δ را می‌توان به عنوان حاصلضرب کرونه‌کر ماتریس دقت اثر زمانی و ماتریس دقت اثر فضایی نوشت.

۲.۱.۶ اثرات اصلی

ماتریس ساختار برای اثر زمانی، ν ، را می‌توان به صورت $R_\nu = \tilde{H} - H$ نوشت، که در آن H ماتریس ساختار همسایگی اثر زمانی و $\tilde{H}$ یک ماتریس قطری با عناصر روی قطر برابر با جمع ردیف ماتریس H است. در اینجا، مدل قدم زدن تصادفی مرتبه اول را برای مدل‌بندی اثر زمانی در نظر می‌گیریم که ماتریس دقت آن برابر است با

$$R_\nu = \begin{bmatrix} 1 & -1 & & & \\ 1 & 2 & -1 & & \\ & \ddots & \ddots & \ddots & \\ & & -1 & 2 & -1 \\ & & & -1 & 1 \end{bmatrix} = \begin{bmatrix} 1 & & & & \\ & 2 & & & \\ & & \ddots & & \\ & & & 2 & \\ & & & & 1 \end{bmatrix} - \begin{bmatrix} & 1 & & & \\ & & 1 & & \\ & & & \ddots & \\ & & & & 1 & \\ & & & & & 1 \end{bmatrix}.$$

به‌طور مشابه، ماتریس دقت برای اثر فضایی را نیز می‌توان به صورت $R_u = \tilde{G} - G$ نوشت که در آن G ماتریس ساختار همسایگی فضایی با عناصر

$$G_{ij} = \begin{cases} 1 & j \sim i \\ 0 & \text{بقیه جاها} \end{cases}$$

است و $\tilde{G}$ یک ماتریس قطری با عناصر روی قطر برابر با جمع ردیف ماتریس G است.

۳.۱.۶ اثرات متقابل فضایی-زمانی

در تحلیل داده‌های با اثر فضایی شبکه‌ای بر اساس کلیتون و همکاران (۱۹۹۶)، نور-هلد (۲۰۰۰) و بلانجیاردو و همکاران (۲۰۱۵)، ماتریس کوواریانس فضایی-زمانی را طوری در نظر می‌گیریم که جداپذیر باشد و از حاصلضرب کرونه‌کر ماتریس دقت اثر فضایی و ماتریس دقت اثر زمانی به دست آید. در حالت کلی، چهار نوع از این اثر به صورت زیر قابل تعریف هستند.

۱. **اثر فضایی-زمانی نوع اول:** اگر دو اثر فضایی و زمانی غیرساختاری باشند، آن‌گاه

$$R_\delta = R_u \otimes R_\nu = I \otimes I = I.$$

۲. **اثر فضایی-زمانی نوع دوم:** در این حالت، اثر فضایی غیرساختاری را با اثر زمانی ساختاری ترکیب می‌کنند. بنابراین، ماتریس ساختاری اثر فضایی-زمانی به صورت

$$R_\delta = R_u \otimes R_\nu = I \otimes R_\nu,$$

بازنویسی می‌شود. در این نوع، ساختاری در فضا وجود ندارد اما روند زمانی از یک ناحیه به ناحیه دیگر فرق می‌کند.

۳. **اثر فضایی-زمانی نوع سوم:** در این نوع اثر متقابل، ساختار در اثر فضایی وجود دارد اما اثر زمانی غیرساختاری است. پس

$$\mathbf{R}_\delta = \mathbf{R}_\nu \otimes \mathbf{R}_u$$

که در آن $\mathbf{R}_u$ ماتریس ساختار همسایگی مدل بی‌سیگ است و $\mathbf{R}_\nu = \mathbf{I}$. در حقیقت ناهمگنی فضایی از یک نقطه زمانی تا نقطه زمانی دیگر متفاوت است بدون این‌که اثر زمانی ساختاری باشد.

۴. **اثر فضایی نوع چهارم:** در این نوع، هر دو ماتریس دقت اثر زمانی و اثر فضایی از نوع ساختاری فرض می‌شوند. بنابراین، ماتریس ساختاری اثر فضایی-زمانی برابر است با $\mathbf{R}_\delta = \mathbf{R}_u \otimes \mathbf{R}_\nu$. نوع چهارم پیچیده‌ترین و پرکاربردترین نوع اثر متقابل فضایی-زمانی است. ماتریس دقت در این نوع به صورت

$$\begin{aligned} \tau_\delta \mathbf{R}_\nu \otimes \mathbf{R}_u &= \tau_\delta (\tilde{\mathbf{H}} - \mathbf{H}) \otimes (\tilde{\mathbf{G}} - \mathbf{G}) \\ &= \tau_\delta (\tilde{\mathbf{H}} \otimes \tilde{\mathbf{G}} - \tilde{\mathbf{H}} \otimes \mathbf{G} - \mathbf{H} \otimes \tilde{\mathbf{G}} + \mathbf{H} \otimes \mathbf{G}), \end{aligned} \quad (2.6)$$

به‌دست می‌آید. اولین جمله (۲.۶) یک ماتریس قطری است که مولفه‌هایش از ضرب تعداد همسایه‌های فضایی در تعداد همسایه‌های زمانی نتیجه می‌شوند. جمله دوم شامل همسایه‌های فضایی هستند که با تعداد همسایه‌های زمانی وزن دار شده‌اند. همچنین، جمله سوم همسایه‌های زمانی هستند که با تعداد همسایه‌های فضایی وزن دار شده‌اند. آخرین جمله هم ساختار همسایگی فضایی-زمانی است.

برای پیاده کردن اثر فضایی-زمانی، بنا بر نور-هلد (۲۰۰۰) باید قید مجموع برابر با صفر را برای این اثر نیز در نظر بگیریم. زیرا مدل ذاتی است و برای شناسایی نیاز به قید دارد. برای شرح بیشتر بخش زیر را دنبال کنید.

۴.۱.۶ قیدها و فضای صفر

از توجیح قیدهای مربوط به اثرات اصلی شروع می‌کنیم و به قیدهای مربوط به اثرات متقابل نوع دوم، سوم و چهارم می‌رسیم. بر اساس رو و هلد (۲۰۰۵)، ماتریس دقت یک IGMRF مرتبه اول را با ضرب پارامتر دقت در ماتریس ساختاری نشان دادیم. این ماتریس پرتبه نیست و کمبود رتبه‌ای برابر با یک دارد؛ به این معنی که فضای صفر در آن یک بعدی است. در نتیجه، ماتریس دقت آن یک مقدار ویژه برابر با صفر دارد. همه عناصر بردار ویژه متناظر هم مقدار یکسانی دارند. بنا بر رو و هلد (۲۰۰۵) (فصل سوم) مقداری که عناصر بردار ویژه به

خود اختصاص می‌دهد، می‌تواند یک باشد. بنابراین در این جا بردار $\frac{1}{n}\mathbf{1}$ را در نظر می‌گیریم، که در آن منظور از $\mathbf{1}$ برداری از 1 هاست. همان گونه که در فصل چهارم نیز اشاره شد، یک راه حل کابردی این است که مقدار کوچکی به عناصر ماتریس دقت فرایند IGMRF اضافه کنیم بدون این که چگالی آن عوض شود. راه دیگر این است که به کمک قید مجموع برابر با صفر، فرایند را شناسایی پذیر کنیم.

برای اثر متقابل نوع دوم، فرایند زمانی در δ یک IGMRF مرتبه اول مانند ν است. ماتریس ساختاری نیز $\mathbf{I} \otimes \mathbf{R}_\nu$ است، که در آن $\mathbf{I}$ یک ماتریس همبستگی با بعد n است. بنابراین، همه کمبود رتبه برابر با n است. فضای صفر برابر است با ضرب کروشه فضای صفر ماتریس دقت اثر زمانی در ماتریس همبستگی با بعد n ؛ یعنی

$$\mathbf{W}_2 = \frac{1}{n} \mathbf{1}_{T \times 1} \otimes \mathbf{I}_n.$$

بنابراین به n قید مجموع برابر با صفر، هر کدام برای یک ناحیه یا هر کدام از n فرایند زمانی نیاز داریم.

اثر متقابل نوع سوم به هر نقطه زمانی t یک میدان فضایی اختصاص می‌دهد. بنابراین δ یک IGMRF مرتبه اول فضایی مانند $\mathbf{u}$ است. کمبود رتبه $\mathbf{I}_T \otimes \mathbf{R}_u$ نیز برابر T است و فضای صفر،

$$\mathbf{W}_3 = \mathbf{I}_T \otimes \frac{1}{T} \mathbf{1}_{n \times 1}$$

است. بنابراین، قید مجموع برابر با صفر برای هر کدام از T میدان فضایی است. اثر متقابل نوع چهارم یک IGMRF در دو جهت فضایی و زمانی است. کمبود رتبه در این حالت برابر با $n + T - 1$ است؛ زیرا به بیان ساده، برای n فرایند زمانی نیاز به قید مجموع برابر با صفر داریم و همین طور برای میدان فضایی در هر یک از $T - 1$ نقطه زمانی. به طریقی دیگر، می‌توان گفت $T - 1$ قید برای اثر نوع سوم و n تا برای اثر نوع دوم نیاز داریم که در مجموع برابر است با $n + T - 1$. برای وارد کردن اثر فضایی-زمانی در مدل چندجمله‌ای و سپس برازش آن به کمک روش INLA، دوباره از تبدیل MP استفاده می‌کنیم. رده‌ها را طوری در نظر می‌گیریم که پاسخ همانند فصل چهارم دارای فراوانی باشد؛ چرا که در تحلیل داده‌های فضایی-زمانی با داده‌های چندجمله‌ای روبرو هستیم که اغلب از فراوانی برخوردارند. راه حل ما برای برطرف شدن مشکل شناسایی پذیری در حالت فضایی-زمانی، اضافه کردن قید مجموع برابر با صفر به مدل است. در این صورت، قیدهای زیادی در مدل داریم که توانایی نسبتاً محدود INLA در به کارگیری قیدها ما را با چالش روبرو می‌کند.

در روش INLA می‌توان قیدهای خطی به صورت $A\chi = e$ را در مدل به کار گرفت، که در آن χ بردار مولفه‌های پنهان با یک میدان گوسی است. طول بردار e و ردیف ماتریس A برابر با طول قیدهای مورد نیاز است. هزینه مثلاً k قید در INLA برابر است با $O(nk^2)$ که در آن n در آن اندازه نمونه و در محاسبات ما برابر با اندازه نقاط فضایی-زمانی است. این هزینه در هر مرحله از الگوریتم INLA باید تکرار شود. بدیهی است هرچه تعداد قیدها افزایش یابد، پیچیدگی و

هزینه محاسباتی بالاتر خواهد رفت. بنابراین، متأسفانه برای به‌کارگیری قیدها در INLA با مشکل مواجه هستیم. در حالت کلی دو راه برای دوری از این مشکل وجود دارد:

۱. پارامتربندی متفاوت

۲. پیاده کردن قیدها بر روی نمونه‌های توزیع پسین.

۲.۶ راه‌های دوری از قید مجموع مساوی با صفر در اثرات متقابل

در این بخش راه حل موجود و راه حل پیشنهادی خود را برای رفع این مشکل بیان می‌کنیم. راه حل اول استفاده از یک نوع پارامتربندی متفاوت است که اولین بار توسط کلیتون (۱۹۹۶) معرفی شد. راه حل دوم، در نظر گرفتن قید بعد از برازش مدل است. در زیر جزئیات بیشتری از این دو روش را تشریح می‌کنیم.

۱.۲.۶ پارامتربندی متفاوت

این روش معمولاً در رگرسیون خطی زمانی که با متغیرهای عاملی روبرو باشیم، به کار می‌رود. در این روش یک سطح به‌عنوان پایه در نظر گرفته می‌شود و بقیه سطوح با این سطح پایه مقایسه و تفسیر می‌شوند. کلیتون (۱۹۹۶) آن را برای دوری از ذاتی بودن IGMRF پیشنهاد کرد.

فرض کنید x یک بردار تصادفی از یک IGMRF مرتبه اول با ماتریس دقت Q_x و کمبود رتبه k است. به دنبال استفاده از پارامتربندی متفاوت، k عنصر از بردار x برابر صفر خواهند شد. فرض کنید I یک بردار به طول k برای شناسایی عناصری از x که صفر شدند، است. وقتی که عنصری از x صفر می‌شود، آن‌گاه ماتریس دقت توزیع شرطی عناصر باقیمانده تحت شرط $x|_{xI=0}$ همان ماتریس دقت قبلی است بدون سطر I و ستون I . حال فرض کنید $x^{(I)}$ یک بردار هم طول با بردار x باشد که تعداد I عنصر ورودی از آن برابر صفر است. ماتریس دقت عناصر باقیمانده برابر با ماتریس دقت $x|_{xI=0}$ است. فضای صفر Q_x یعنی N_x را می‌توان به‌عنوان ترکیب خطی در محاسبه کلی قید به صورت $x|_{N_x x=0}$ به کار گرفت. در ادامه یک کاربرد از این ایده را برای مدل قدم زدن تصادفی مرتبه اول نشان می‌دهیم.

پارامتربندی متفاوت در مدل قدم زدن تصادفی

مدل قدم زدن تصادفی مرتبه اول را با بردار $\nu = (\nu_1, \dots, \nu_T)'$ و ماتریس دقت $\tau_\nu R_\nu$ در نظر بگیرید. به جای استفاده از قید مجموع مساوی با صفر یعنی $\nu'_1 = 0$ ، یکی از عناصر ν ، مثلاً

ν_l را مساوی صفر قرار می‌دهیم. توزیع شرطی $\nu_l | \nu_l = 0$ ، یک توزیع گوسی است که ماتریس دقت آن همان ماتریس دقت فرآیند قدم زدن تصادفی مرتبه اول اصلی است که ردیف l و ستون l آن حذف شده باشد (کلیتون، ۱۹۹۶). پیامد این کار این است که مدل پارامتربندی شده $\nu^{(l)} = (\nu_1^{(l)}, \dots, \nu_T^{(l)})'$ ، تفاوت‌های $\nu_1 - \nu_l, \dots, \nu_T - \nu_l$ را مدل‌بندی می‌کند. در حقیقت اثرات اصلی با ν_l به $\nu_1^{(l)}, \dots, \nu_{l-1}^{(l)}, \nu_{l+1}^{(l)}, \dots, \nu_T^{(l)}$ منتقل می‌شوند، در حالی که ν_l خودش نیست. با این وجود، می‌توان اثرات اصلی ν_t را به‌وسیله اثرات منتقل شده $\nu_t^{(l)}$ به صورت

$$\nu_t = \nu_t^{(l)} - \frac{1}{T} \sum_{t=1}^T \nu_t^{(l)}$$

بازیابی کرد، که در واقع ترکیب خطی تعریف شده روی $\nu^{(l)}$ است.

پارامتربندی متفاوت در اثرات متقابل

برای این که بتوانیم روش پارامتربندی متفاوت را برای اثرات متقابل فضایی-زمانی پیاده کنیم، باید یک نقطه زمانی l برای اثر متقابل نوع دوم و یک نقطه فضایی a برای اثر متقابل نوع سوم و همزمان یک نقطه زمانی l و یک نقطه فضایی a را برای اثر متقابل نوع چهارم به‌عنوان سطح پایه انتخاب کنیم. بنابراین ترکیبات خطی برای اثر متقابل نوع دوم همانند حالت اثر زمانی به صورت

$$\delta_{it} = \delta_{it}^{(l)} - \frac{1}{T} \sum_{t=1}^T \delta_{it}^{(l)},$$

بازنویسی می‌شود. اگر $t = l$ باشد، آن گاه $\delta_{it}^{(l)} = 0$. به‌طور مشابه برای اثر نوع سوم می‌توان نوشت

$$\delta_{it} = \delta_{it}^{(a)} - \frac{1}{n} \sum_{i=1}^n \delta_{it}^{(a)}.$$

برای اثر نوع چهارم باید هر دو جهت فضا و زمان را در نظر بگیریم. بنابراین

$$\delta_{it} = \delta_{it}^{(a,l)} - \frac{1}{n} \sum_{i=1}^n \delta_{it}^{(a,l)} - \frac{1}{T} \sum_{t=1}^T \delta_{it}^{(a,l)} + \frac{1}{nT} \sum_{i=1}^n \sum_{t=1}^T \delta_{it}^{(a,l)}.$$

توجه داشته باشید که روابط بالا از روی فضای صفر ماتریس دقت اثرات اصلی و بر اساس کلیتون (۱۹۹۶) به‌دست آمده‌اند.

متأسفانه به‌کارگیری این ایده در کار ما همچنان دست و بال ما را بسته نگه داشت، زیرا همچنان با قیدهای زیادی رو به رو هستیم اگر چه نوع ترکیب آن عوض شده باشد. افزون بر این، همه این قیدها در مدل چندجمله‌ای برای سه رده تکرار می‌شود. این یعنی اگر فردی داده‌ای به حجم مثلاً ۱۰۰۰ داشته باشد، در عمل مدل داده‌ای به حجم ۳۰۰۰ یعنی سه برابر را برازش می‌دهد که همه این قیدها در آن تکرار می‌شود. این حجم با افزایش تعداد رده‌ها نیز

افزایش می‌یابد. چیزی که به قول دوست نروژی‌ام، ماریای عزیز، به کابوس شب شبیه است. لازم به ذکر است که اجرای آن با حجم داده شبیه‌سازی شده پایین نیز ما را به نتایج دلخواه نرساند!

۲.۲.۶ پیاده کردن قیدها بر روی نمونه‌های توزیع پسین

در تشریح این ایده فرض کنید Y دارای توزیع گوسی با ماتریس دقت Q است. همچنین فرض کنید A ماتریس تعریف قیدها است. قیدها می‌توانند برای رفع ذاتی بودن مدل یا رفع مشکل شناسایی‌پذیری پارامترهای مدل باشند. با فرض وارون‌پذیری Q ، مراحل زیر را انجام می‌دهیم:

۱. برازش مدل بدون قیدها و تولید نمونه از توزیع پسین.
 ۲. پیاده کردن قیدها بر روی نمونه‌های تولیدشده پسینی.
- ابتدا این روش را بر روی یک مدل زمانی کلاسیک به کار می‌گیریم.

مثال ۱.۲.۶. فرض کنید یک مدل کلاسیک با معادله

$$y_i = \alpha + x_i + \epsilon_i,$$

داریم که در آن x از یک مدل قدم‌زدن تصادفی مرتبه اول پیروی می‌کند. مؤلفه ϵ_i نیز دارای توزیع نرمال با میانگین صفر و دقت متفاوت به ازای هر مشاهده است. برای دوری از ترکیب شدن x با عرض از مبدأ α باید قید $\sum_i x_i = 0$ را در نظر بگیریم. از این مدل، نمونه‌ای به حجم ۴۰ تولید کردیم. دقت جمله ϵ_i را به صورت $\frac{1}{\sqrt{\exp(N(n, 1, 3))}}$ در نظر گرفتیم. مقدار α را نیز برابر با یک گرفتیم. بنابراین ابتدا مدل را با در نظر گرفتن کمبود رتبه در ماتریس دقت و پیشین‌های پیش‌فرض بسته R-INLA برازش دادیم و در خروجی از نمونه‌های پسین حاصل به اندازه ۱۰۰۰ نمونه جدا کردیم. سپس ماتریس A را گونه‌ای طراحی کردیم تا با پیاده کردن آن به قید دلخواه برسیم. خروجی حاصل تا چهار رقم اعشار برای چند نمونه اول در جدول ۱.۶ گزارش شده است. همان‌طور که مشهود است، مقادیر برآوردشده به مقادیر شبیه‌سازی شده خیلی نزدیک هستند. بنابراین، به مدل اصلی برمی‌گردیم و اجرای این روش را در آن پی می‌گیریم.

توجه داشته باشید که اهمیت به کارگیری این روش به‌خاطر اهمیت به کارگیری اثرات متقابل در تحلیل داده‌های چندجمله‌ای است که تاکنون افراد (حتی به کمک روش‌های MCMC) فقط موفق به استفاده از اثر نوع اول شده‌اند که آن هم ساختاری ندارد و از این رو مورد توجه نیست!

جدول ۱.۶: مقایسه میانگین و انحراف معیار برآوردشده با مقادیر شبیه‌سازی‌شده در مثال ۱.۲.۶.

میانگین شبیه‌سازی‌شده	میانگین برآوردشده	انحراف استاندارد شبیه‌سازی‌شده	انحراف استاندارد برآوردشده
۱۶/۷۹۹۵	۱۶/۷۹۹۷	۱/۴۹۷۵	۱/۵۱۵۸
۱۴/۴۸۶۱	۱۴/۴۸۶۲	۱/۴۹۷۶	۱/۵۱۵۸
۱۳/۷۱۱۷	۱۳/۷۱۱۹	۱/۴۹۷۷	۱/۵۱۵۸
۱۲/۶۵۶۳	۱۲/۶۵۶۴	۱/۴۹۷۵	۱/۵۱۵۷
۱۲/۵۸۵۸	۱۲/۵۸۰۸	۱/۲۳۹۱	۱/۲۳۳۳
۱۱/۴۰۳۰	۱۱/۳۹۸۳	۱/۲۴۰۱	۱/۲۳۳۸
۱۰/۵۶۴۳	۱۰/۵۵۹۲	۱/۲۳۹۲	۱/۲۳۳۳
۹/۵۵۶۹	۹/۵۵۱۹	۱/۲۳۹۲	۱/۲۳۳۳
۸/۶۵۷۴	۸/۶۵۲۳	۱/۲۳۹۳	۱/۲۳۳۲
۸/۵۷۰۸	۸/۵۶۹۷	۱/۰۰۰۳	۰/۹۹۹۰
۷/۵۲۲۵	۷/۵۲۱۶	۰/۹۹۹۵	۰/۹۹۸۳
۶/۶۴۵۶	۶/۶۴۴۷	۰/۹۹۹۶	۰/۹۹۸۳
۵/۵۳۳۷	۵/۵۳۲۹	۰/۹۹۹۴	۰/۹۹۸۳
۴/۴۵۲۳	۴/۴۵۱۴	۰/۹۹۹۵	۰/۹۹۸۳
۴/۴۲۶۳	۴/۴۱۰۸	۰/۸۵۰۰	۰/۸۵۸۲
۳/۵۶۶۳	۳/۵۵۱۴	۰/۸۵۰۸	۰/۸۵۸۷
۲/۰۱۵۲	۱/۹۹۹۹	۰/۸۴۹۹	۰/۸۵۸۰
۰/۹۰۲۰	۰/۸۸۶۳	۰/۸۵۰۱	۰/۸۵۸۲
۰/۰۰۴۹	۰/۰۲۰۳	۰/۸۴۹۹	۰/۸۵۸۰
۰/۱۰۰۷	۰/۱۰۱۲	۰/۸۵۴۴	۰/۸۵۰۱
۱/۰۵۲۹	۱/۰۵۳۴	۰/۸۵۴۵	۰/۸۵۰۱
۱/۹۷۶۶	۱/۹۷۷۱	۰/۸۵۴۵	۰/۸۵۰۱
۳/۲۲۶۷	۳/۲۲۷۲	۰/۸۵۵۱	۰/۸۵۰۶
۳/۸۳۱۲	۳/۸۳۱۹	۰/۸۵۴۶	۰/۸۵۰۲
۴/۷۳۵۷	۴/۷۳۳۴	۱/۰۱۶۲	۱/۰۲۳۳

۳.۶ مطالعه شبیه‌سازی

در مطالعه شبیه‌سازی ارایه‌شده در این بخش، داده‌ها را از مدل (۲.۴) با تابع درست‌نمایی (۱.۴) با در نظر گرفتن پیشگوی جمعی

$$\eta_{ijt} = x_{ijt}\beta_j + u_{ijt} + \nu_{ijt} + \delta_{ijt}, \quad i = 1, \dots, n, \quad t = 1, \dots, T, \quad j = 1, \dots, J, \quad (3.6)$$

تولید کردیم، که در آن $T = 5$ ، $n = 15$ و $J = 3$ انتخاب شدند. برای سبکی مدل و تمرکز بیشتر بر روی اثر متقابل، عرض از مبدأ را حذف کردیم. اثر زمان برای سه رده را به ترتیب زیر شبیه‌سازی کردیم.

$$t_1 = \sin(2\pi(1 : T)/5^\circ), \quad t_2 = \log(1 : T), \quad t_3 = ((1 : T)/T)^2,$$

و در فرآیند مدل‌بندی و برازش، اثر زمانی ν را به کمک یک فرآیند قدم زدن تصادفی مرتبه اول مدل کردیم. برای تولید اثر فضایی، u ، نواحی $i = 1, \dots, n$ را بر روی یک شبکه منظم 10×10 تعریف کردیم و تحقق‌های آن را از یک مدل بی‌سیگ به کمک تابع `qsampl` در بسته R-INLA شبیه‌سازی کردیم. اثر متقابل δ را نیز از نوع چهارم شبیه‌سازی کردیم. برای محاسبه ماتریس دقت اثر متقابل از ضرب کرونه کر ماتریس دقت اثر زمانی و ماتریس دقت فضایی مدل بی‌سیگ، استفاده کردیم. متغیر تبیینی x برای همه رده‌ها از توزیع نرمال استاندارد شبیه‌سازی شد و ضرایب β_j نیز از یک توزیع نرمال استاندارد تولید و به گونه‌ای مقیاس‌بندی شدند که جمع آن‌ها برابر صفر شود.

لازم به ذکر است برای جلوگیری از ترکیب نشدن اثرها با هم از مقیاس‌بندی مدل (تعریف‌شده در فصل چهارم) بهره‌مند شدیم. در این جا مقیاس‌بندی پیشین‌های تعریف‌شده در مدل یا به عبارتی همان IGMRF‌ها نه تنها به ما کمک می‌کند تا اثرها با هم ترکیب نشوند و برآوردهای حاصل از پسین برای ابرپارامترها قابل مقایسه باشند، بلکه این امکان را می‌دهد تا بتوانیم برای ابرپارامترها نیز توزیع پیشین در نظر بگیریم. این به آن علت است که پارامتر دقت IGMRF مقیاس‌بندی‌شده، دقت کناری است. بنابراین ریشه دوم معکوسش انحراف استاندارد کناری است که اندازه اثر را در مقیاس پیشگوی خطی اندازه‌گیری می‌کند. برازش مدل نیز بدون در نظر گرفتن قیدها اجرا شد.

برای برازش مدل از توزیع پیشین PC برای ابرپارامترهای مدل استفاده کردیم. ایده اصلی این توزیع پیشین بر پایه جریمه کردن اختلاف یک مدل (ساده) پایه و یک مدل با ساختار پیچیده‌تر است. در این ایده، پارامتر طوری در نظر گرفته می‌شود که انحراف از مدل پایه را اندازه‌گیری و جریمه می‌کند. برای اثر تصادفی، مدل پایه یک خط مستقیم است که واریانس صفر دارد. بنابراین واریانس (انحراف معیار) اختلاف مدل پایه از مدل پیچیده‌تر را اندازه‌گیری می‌کند.

فرض کنید $v > 0$ پارامتری باشد که مدل را کنترل می‌کند. هرچه v بزرگتر باشد، مدل پیچیده‌تر است. ایده پیشین PC، تعریف توزیع‌هایی است که در آن‌ها $P(v > q) = \alpha$ که در آن مقادیر q و α توسط کاربر به گونه‌ای تعیین می‌شوند که باور پیشینی او را بر توزیع پیشین القا کنند. در این مثال، به‌طور تجربی، مقادیر مختلف را بررسی کردیم و در نهایت $q = 0.5$ و $\alpha = 0.5$ را انتخاب کردیم. تفسیر این مقادیر به سادگی قابل درک هستند: به‌عنوان نمونه توزیع پیشین PC برای پارامتر دقت مدل بی‌سیگ با این مقادیر منتخب به این معنی است که انتظار داریم ۵۰٪ تغییرات با ساختار فضایی توضیح داده بشود.

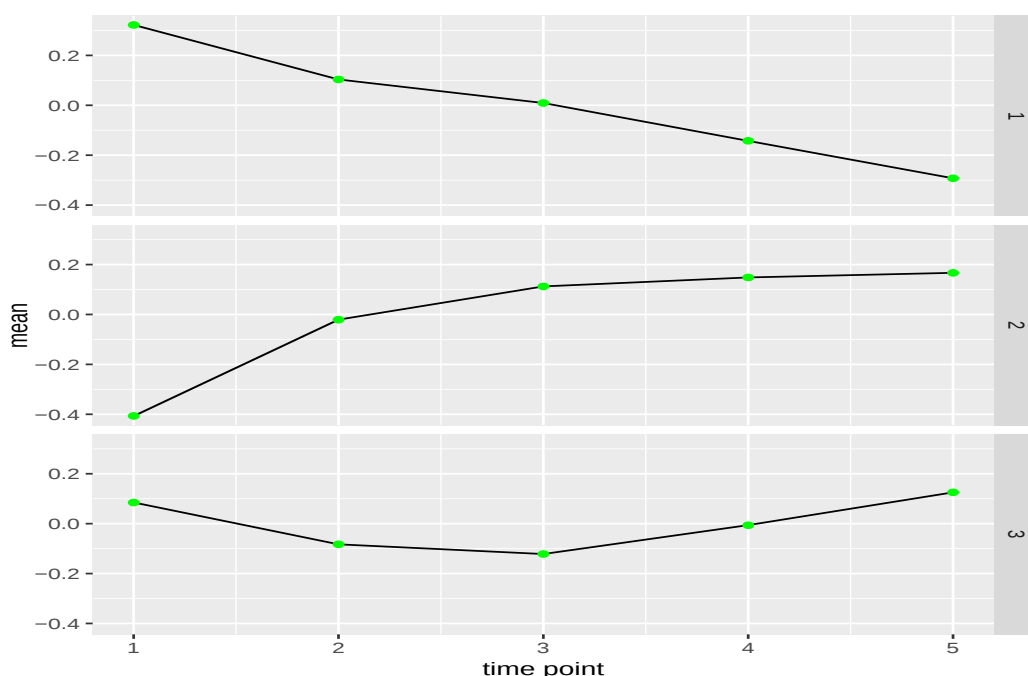
نتایج حاصل از برازش مدل در جدول ۲.۶ و شکل‌های ۱.۶، ۲.۶ و ۳.۶ نمایش داده شده‌اند. با توجه به نتایج می‌توان گفت برآوردهای پسینی هم ضرایب رگرسیونی و هم اثرات فضایی، زمانی و متقابل آن‌ها از دقت قابل قبول برخوردارند. مطالعه شبیه‌سازی حاوی نتایج راضی‌کننده‌ای است. با این حال، نرسیدن به نتایج نظری تاییدکننده ما را مجبور کرد، تا زمان یافتن مبانی نظری برای اثبات ریاضی آن، از روش دیگری برای ادامه کار استفاده کنیم.

برازش مدل با اثر فضایی-زمانی با قید رده گوشه ۱۰۷

از آن جایی که با روش‌های دیگر در مدل فضایی-زمانی به نتیجه دلخواه نرسیدیم به ناچار، در ادامه، مدل را با قید رده گوشه برازش می‌دهیم. اگرچه که در این حالت برآورد مدل از دقت پایین‌تری برخوردار است، اما در حال حاضر تنها راه برازش مدل فضایی-زمانی چندجمله‌ای به شمار می‌رود.

جدول ۲.۶: برآورد میانگین (انحراف استاندارد) پسینی اثرات رگرسیونی ثابت.

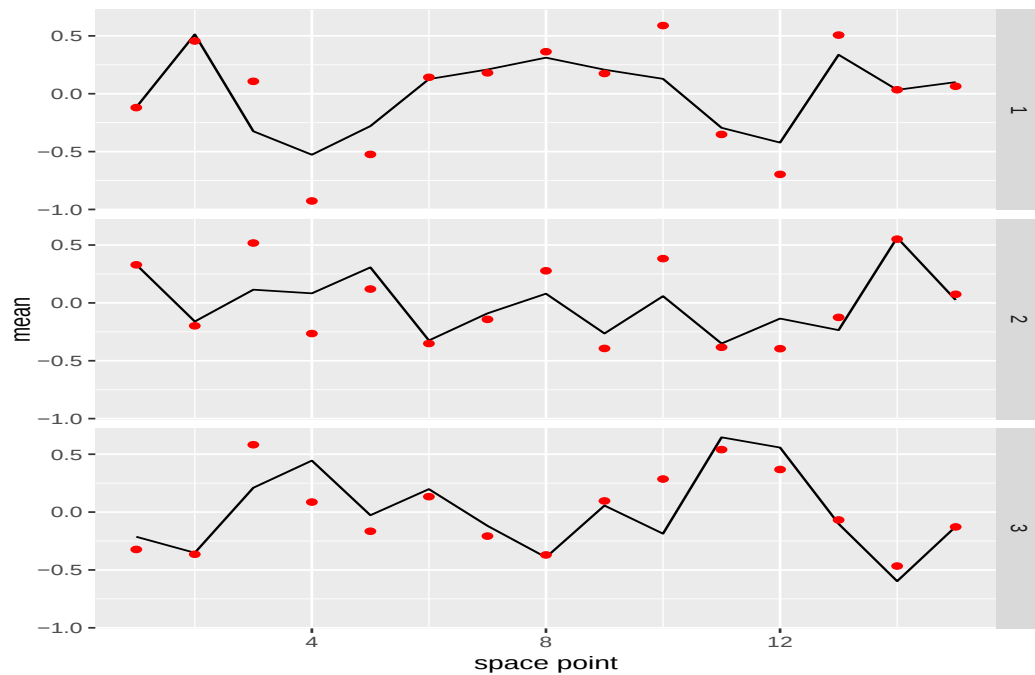
پارامتر	مقدار واقعی	برآورد
β_1	-۰/۱۷۳	-۰/۱۲۱ (۰/۰۰۶)
β_2	۱/۱۷۳	۱/۵۳۱ (۰/۰۰۳)
β_3	-۱/۰۰	-۱/۱۰۸ (۰/۰۰۹)



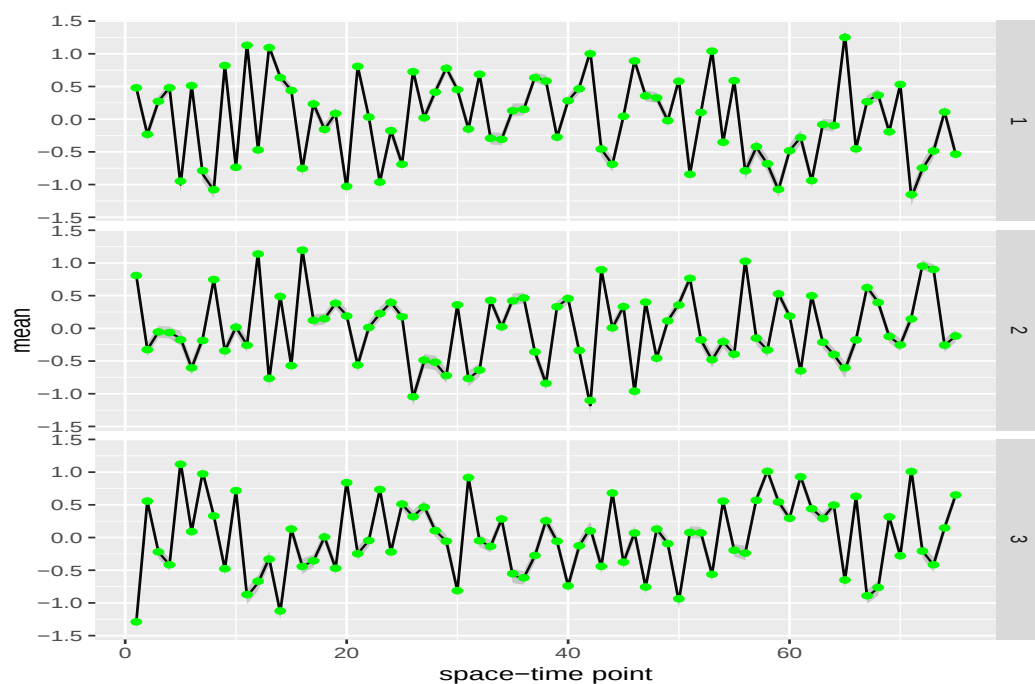
شکل ۱.۶: میانگین پسین اثر زمانی (دایره‌های سبزرنگ) به همراه اثر زمانی شبیه‌سازی شده واقعی (خط مشکی) برای هر سه رده.

۴.۶ برازش مدل با اثر فضایی-زمانی با قید رده گوشه

در این بخش مدل چندجمله‌ای جمعی-ساختاری با اثر فضایی-زمانی را با در نظر گرفتن قید رده گوشه برای شناسایی‌پذیری مطالعه می‌کنیم. اگرچه نتایج مربوط به همه اثرات متقابل خوب بودند، اما به دلیل محدودیت فضای رساله از آوردن همه نتایج خودداری کرده و فقط نتایج مربوط به اثر متقابل نوع چهارم را ذکر می‌کنیم که در باور بسیاری از کاربران علوم دیگر مهمتر و پرثمرتر جلوه کرده است.



شکل ۲.۶: میانگین پسین اثر فضایی (دایره‌های قرمز رنگ) به همراه اثر فضایی شبیه‌سازی شده واقعی (خط مشکی) برای هر سه رده.



شکل ۳.۶: میانگین پسین اثر متقابل فضایی-زمانی (دایره‌های سبز رنگ) به همراه اثر متقابل شبیه‌سازی شده واقعی (خط مشکی) برای هر سه رده.

۱.۴.۶ مطالعه شبیه‌سازی

در این مثال، داده‌ها را بر اساس پیشگوی

$$\eta_{ijt} = \alpha_j + x_{ijt}\beta_j + u_{ijt} + v_{ijt} + \delta_{ijt}, \quad i = 1, \dots, n, \quad t = 1, \dots, T, \quad j = 1, \dots, J$$

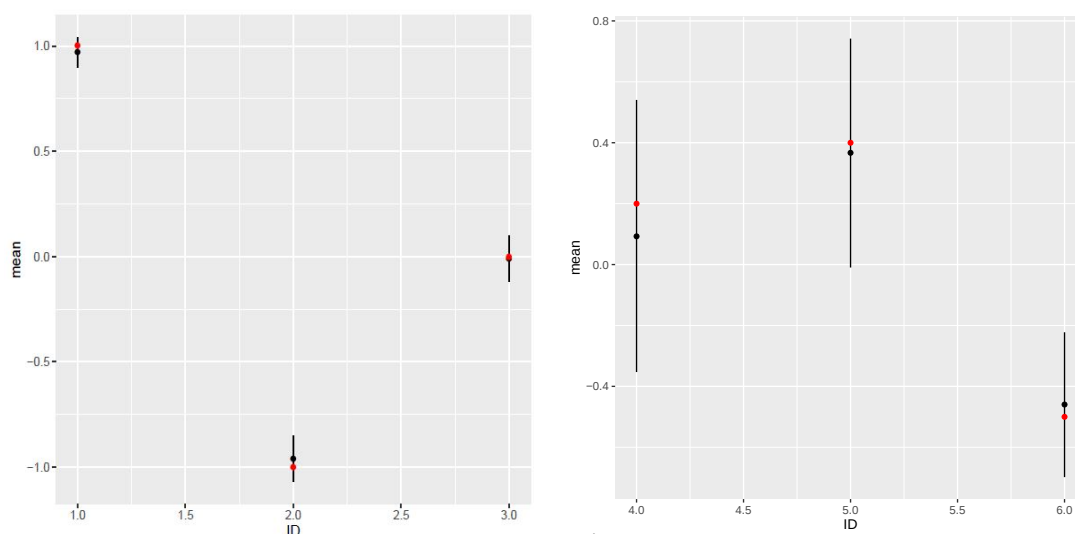
با $n = 30$ ناحیه فضایی، $T = 5$ نقطه زمانی و $J = 3$ رده شبیه‌سازی کردیم. در این شبیه‌سازی، رده اول را به‌عنوان رده گوشه در نظر گرفتیم. بنابراین قرار دادیم

$$\alpha_1 = 0, u_{i \setminus t} = 0, \delta_{i \setminus t} = 0, \nu_{i \setminus t} = 0, \beta_1 = 0.$$

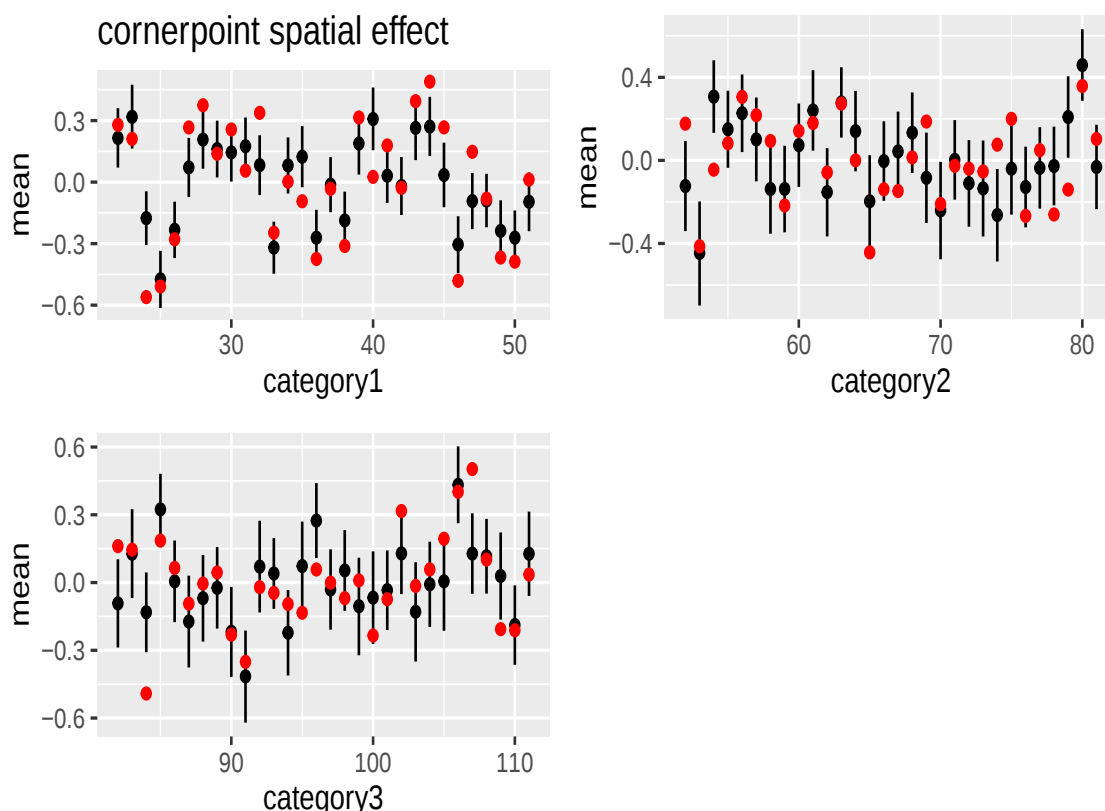
متغیر تبیینی و پارامترهای عرض از مبدأ و ضرایب رگرسیونی را از توزیع نرمال استاندارد تولید کردیم. مشابه با مثال شبیه‌سازی قبل از مقیاس‌بندی مدل و توزیع پیشین PC برای پارامترهای دقت اثرات پنهان استفاده کردیم. اثر زمانی را به کمک توابع

$$t_1 = \sin(t/0.3), t_2 = \cos(t/0.6), t_3 = \cos(t/0.4),$$

تولید کردیم. سایر جزئیات مشابه مثال قبل است. نتایج برازش مدل برای پارامترهای رگرسیونی در شکل ۴.۶ و برآوردهای اثرات فضایی، زمانی و فضایی-زمانی به ترتیب در شکل‌های ۵.۶، ۶.۶ و ۷.۶ نمایش داده شده‌اند. در این شکل‌ها به ترتیب از چپ به راست مربوط به رده‌های اول تا سوم است. پسین کناری مربوط به دقت‌ها نیز در شکل ۱.۴.۶ و احتمال پسین برآوردشده هر رده برای ۱۰۰ نقطه فضایی در شکل ۹.۶ رسم شده‌اند. لازم به ذکر است قید $\sum_i u_{ij} = 0$ نیز به کار گرفته شد تا مدل بی‌سیگ قابل شناسایی باشد. نتایج مربوط به برآوردهای هر سه نوع اثر فضایی، زمانی و فضایی-زمانی در اکثر نواحی، تفاوت نسبتاً کمی را بین مقادیر برآوردشده و مقدار واقعی نشان می‌دهد. برآورد میانگین پسین مربوط به اثر زمان بهتر از برآورد مربوط به سایر اثرها است. احتمال‌های برآوردشده در هر رده، مقادیر نزدیک به مقادیر شبیه‌سازی‌شده را نشان می‌دهد. اما چگالی‌های کناری پسین پارامترهای دقت به‌غیر از اثر زمانی، نادقیق هستند و کم‌برآوردی را نشان می‌دهند.



شکل ۴.۶: برآوردهای میانگین پسین اثر عرض از مبدأ با فاصله اطمینان ۹۵٪ (چپ)، برآوردهای میانگین پسین اثر ضریب رگرسیونی با فاصله اطمینان ۹۵٪ (راست). برآوردها به رنگ قرمز نمایش داده شده‌اند و مقادیر واقعی به رنگ مشکی.



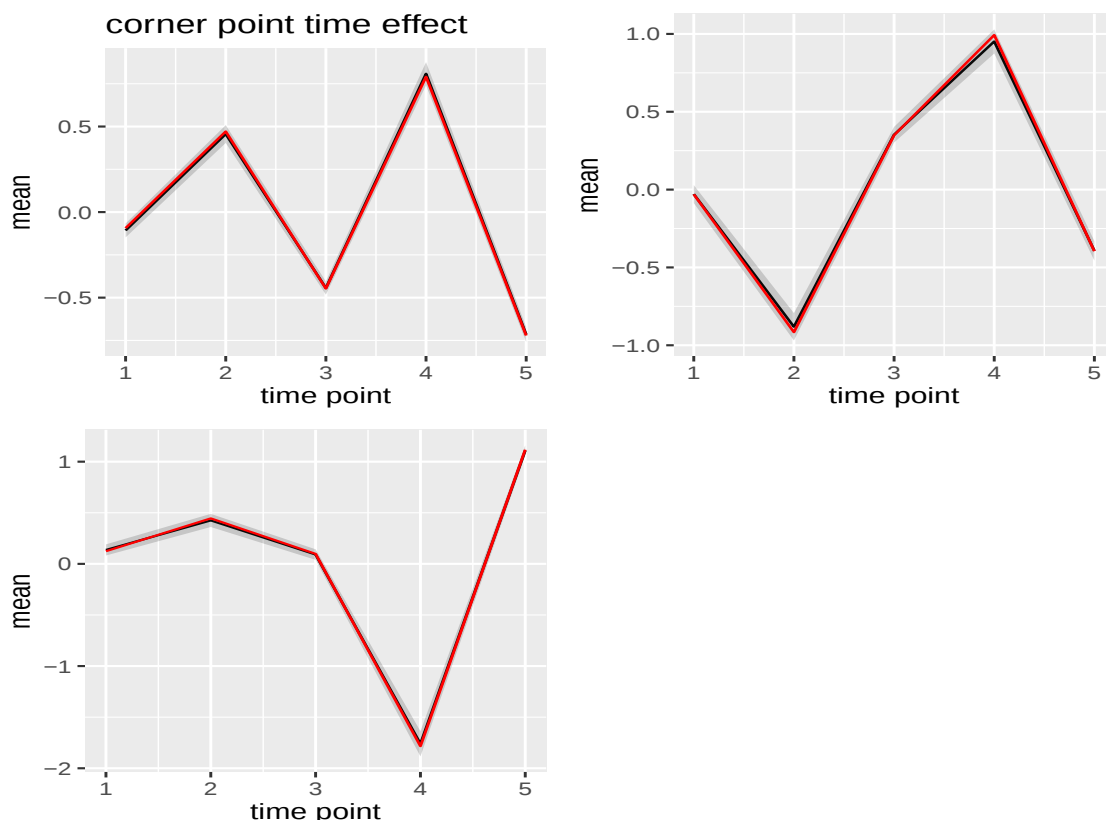
شکل ۵.۶: برآوردهای میانگین پسین اثر فضایی در هر رده با فاصله اطمینان ۹۵٪. رده‌ها در شکل از سمت چپ به ترتیب اول تا سوم هستند. مقادیر واقعی به رنگ مشکی و برآوردها به رنگ قرمز هستند.

۵.۶ خلاصه فصل و نتیجه‌گیری

می‌دانیم با پیشرفت در فن‌آوری ثبت و ذخیره‌سازی داده‌ها و دسترسی به داده‌هایی که هم‌زمان در زمان و مکان دچار تحول می‌شوند و فرآیند تغییر را طی می‌کنند، بسیار آسانتر از گذشته است. همین امر موجب شده است تا افراد در پی کشف فرآیند موجود در داده‌ها باشند و مدل‌های مختلف و روش‌های آماری را به کار گیرند. در این میان داده‌های چندجمله‌ای که در ذات خود طبیعت هستند، از این امر مستثنی نیستند و توجه به آن‌ها بیشتر از قبل ضروری است. در این فصل مدل چندجمله‌ای جمعی-ساختاری با اثر فضایی-زمانی را مورد بررسی و مطالعه قرار دادیم. مشکلات پیش روی شناسایی‌پذیری و راه حل‌ها را بحث و ارزیابی کردیم. سرانجام قید رده گوشه را به عنوان توجیه‌پذیرترین راه حل موجود که پشتوانه نظری ریاضی دارد، پذیرفتیم و عملکرد قابل قبول آن را در یک مطالعه شبیه‌سازی نشان دادیم.

۶.۶ آینده پژوهش

با توجه به مطالعاتی که در طول سال‌های گذشته در باب شناسایی‌پذیری و مدل‌بندی داده‌های رسته‌ای به‌ویژه از نوع چندجمله‌ای انجام داده‌ایم و مشکلات موجود را شناسایی کردیم، در



شکل ۶.۶: برآوردهای میانگین پسین اثر زمانی در هر رده با فاصله اطمینان ۹۵٪. رده‌ها در شکل از سمت چپ به ترتیب اول تا سوم هستند. مقادیر واقعی به رنگ مشکی و برآوردها به رنگ قرمز هستند.

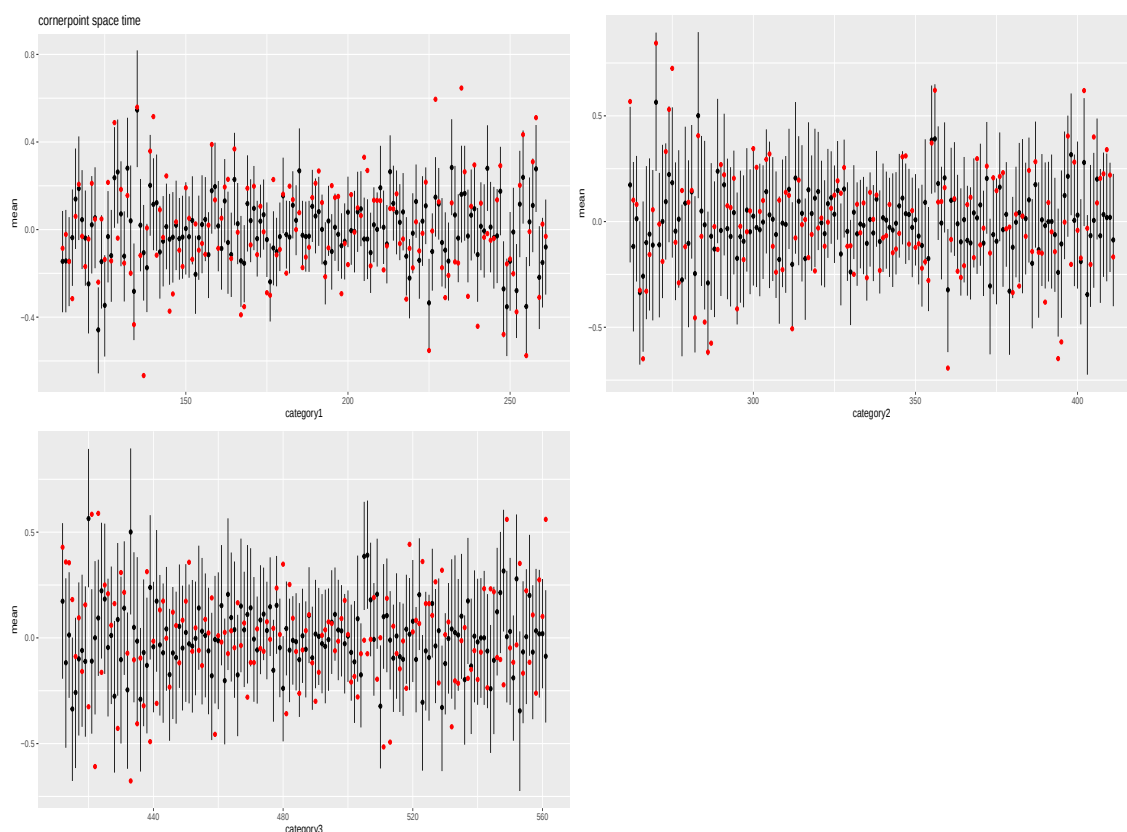
حالت کلی سه راه را برای ادامه کار به شرح زیر بررسی می‌توان کرد:

۱. **عوض کردن نوع قید** به جای بررسی قید به صورت مجموع مساوی با صفر از حالت کلی‌تر $Ax \leq b$ استفاده نماییم. نوع دیگر از قیدها وجود دارد که می‌تواند بر روی نمونه‌ها و احتمال‌ها تعریف شود. مثلاً بررسی قید $(2J - 1) \geq n_i$ یعنی مجموع آزمایش‌ها برای هر فرد بیشتر از دو برابر تفاضل تعداد کل رده‌ها از یک باشد.

۲. **عوض کردن رهیافت به کارگیری مدل** همان‌گونه که قبلاً نیز گفته شد، INLA در به کارگیری قیدهایی غیر خطی ناتوان است و در به کارگیری قیدها خطی توانایی محدودی دارد. بنابراین باید به دنبال روش دیگری باشیم که کارایی برابر با روش INLA داشته باشد و یا حتی از آن کاراتر هم باشد. این روش می‌تواند استنباط تغییراتی^۱ (بلی و همکاران، ۲۰۱۷) با توسعه بیزی و وارد کردن اثر فضایی و فضایی-زمانی باشد که باید در ابتدا آن را به کار بگیریم و کارایی آن را با INLA مورد مقایسه قرار دهیم که قطعاً پیاده کردن آن به سادگی گفتن آن نیست.

۳. **تغییر مدل** به جای استفاده از مدل‌های جمعی-ساختاری چندجمله‌ای از مدل‌های جمعی-ساختاری رگرسیون درختی بیزی استفاده کنیم که یک LGM است و برای

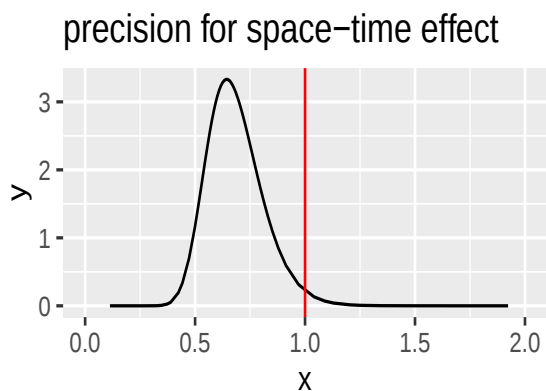
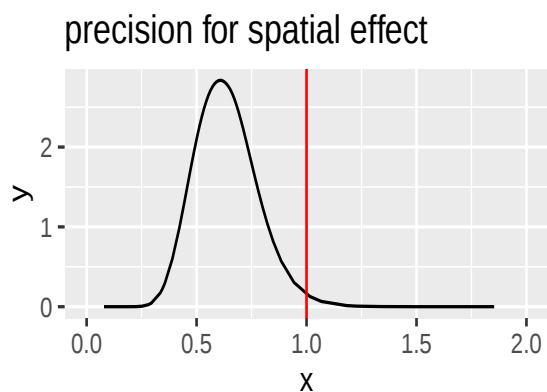
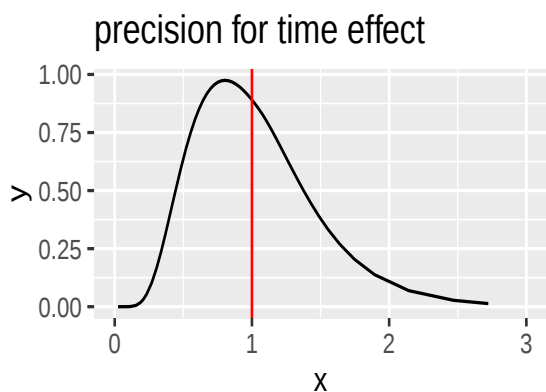
^۱Variational inference



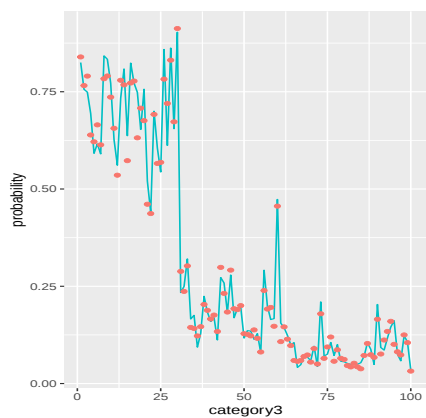
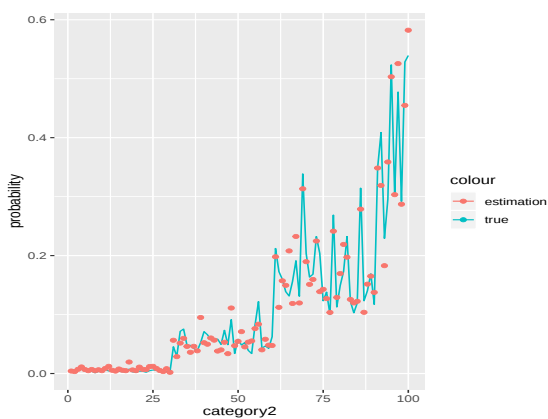
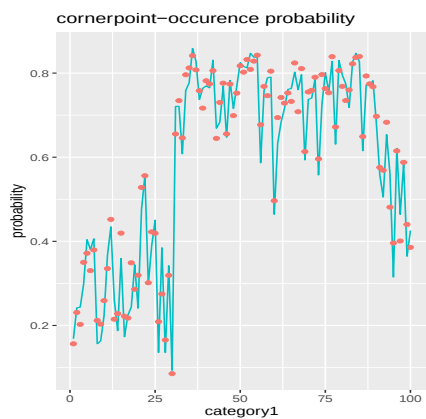
شکل ۷.۶: برآوردهای میانگین پسین اثر فضایی-زمانی در هر رده با فاصله اطمینان ۹۵٪. رده‌ها در شکل از سمت چپ به ترتیب اول تا سوم هستند مقادیر واقعی به رنگ مشکی و برآوردها به رنگ قرمز هستند.

شناسایی نیاز به قید ندارد. بنابراین همچنان قادر خواهیم بود تا از روش کارای INLA استفاده کنیم که البته پیاده‌کردن آن به سادگی گفتن آن نیست و نیاز به زمان زیاد و پیداکردن روشی برای واردکردن وابستگی‌های خاص به مدل دارد.

روش‌های بیان شده در بالا که در طول تحقیقات خود در یک سال اخیر به آن‌ها دست‌یافته‌ایم را می‌توان با صرف مدت زمان بیشتر به کار گرفت و نتایج آن را در صورت قبولی ارائه داد تا این‌گونه کمکی به همه کاربران مدل چندجمله‌ای به‌خصوص در رشته‌های غیر آمار که تعداد آن‌ها کم نیست، کرده باشیم. از طرفی از مشکلات مربوط به قید رده گوشه نیز دوری کرده‌ایم.



شکل ۸.۶: برآوردهای چگالی‌های کناری دقت برای هر اثر.



شکل ۹.۶: برآورد احتمال پسین مربوط به هر رده برای ۱۰۰ نقطه فضایی-زمانی. رده‌ها در شکل از سمت چپ به ترتیب اول تا سوم هستند.

مراجع

- [۱] محمدزاده م. (۱۳۹۱)، ”آمار فضایی و کاربردهای آن”، چاپ اول، مرکز نشر آثار علمی دانشگاه تربیت مدرس، تهران.
- [2] Abramowitz, M. and Stegun, I. A. (1972). "Bessel functions J and Y". in Handbook of Mathematical Functions with Formulas, Graphs and Mathematical Tables, 9th printing. **New York: Dover**. 358-364.
- [3] Agresti, A. (1993). "Computing conditional maximum likelihood estimates for generalized Rasch models using simple loglinear models with diagonals parameters". **Scandinavian Journal of Statistic**. 20, 63-71.
- [4] Agresti, A. (2019). "Introduction to Categorical Data Analysis". **Wiley, New Jersey**. USA.
- [5] Agresti, A. (2002). "Categorical data analysis". second Edition. **New York: Wiley**. USA.
- [6] Agresti, A. (2013). "Categorical data analysis". 3d Edition. **New York: Wiley**. USA.
- [7] Alvarez, R.M. and Nagler, J., (1995). "Economics, issue and the Perot candidacy: voter choice in the 1992 presidential election". **American Journal of Political Science**. 39, 714-744.
- [8] Alvarez, R.M. and Nagler, J., (1998). "When politics and models collide: estimating models of multiparty competition". **American Journal of Political Science**. 42, 55-96.
- [9] Alvarez, R.M. and Nagler, J., (2001). "Correlated disturbances in discrete choice models: a comparison of multinomial probit and logit models". **Political Analysis**. Forthcoming. 914.

-
- [10] Alvarez, R.M. and Nagler, J., Bowler, S., (2000). "Issues, economics, and the dynamics of multiparty elections: the British 1987 general election". **American Political Science Review**. 94, 131–149.
 - [11] Ambo, B.T., Ma, J. and Fu, C. (2020). "Investigating influence factors of traffic violation using multinomial logit method". **International Journal of Injuries Control Safety Promotion**. 28(1), 78-85.
 - [12] Bakka , H., Rue., H., Arne Fuglstad, G., Riebler., A., Bolin, D., Illian., J., Krainski., E., Simpson., D. and Lindgren., F. (2018). "Spatial modelling with R-INLA: A review". **computational statistics**. 1400-1443.
 - [13] Baker, S.G. (1994). "The multinomial-Poisson transformation". **Journal of Royal Statistical Society**. Ser D (The Statistician). 43(4), 495-504.
 - [14] Bansal, P., Krueger, R., Bierlaire, M., Daziano, R.A. and Rashidi, T.H. (2020). "Bayesian estimation of mixed multinomial logit models: advances and simulation-based evaluations". **ransportation Research Part B: Methodological**. 131, 124–142.
 - [15] Barmi, H. E. and Dykstra. R. L. (1998). "Maximum likelihood estimates via duality for log convex models when cell probabilities are subject to convex constraints". **The Annals of Statistics**. 26, 1878-1893.
 - [16] Bartlett, M. S. (1937). "Some examples of statistical methods of research in agriculture and applied biology". **Journal of Royal Statistics Society Supplementary**. 4, 137183.
 - [17] Basse, G. and Bojinov, I. (2020). "A general theory of identification". **Harvard Business School**. 20-86.
 - [18] Basu, P.A. (1983). "Identifiability". In **Encyclopedia of Statistical Sciences** (S. Kotz and N. L. Johnson, eds.), Wiley Interscience. 4, 2–6.
 - [19] Bhattacharya, L., Li, L. and Victor Patrangenaru. (2016). **A Course in Mathematical Statistics and Large Sample Theory**. Springer-Verlag New York.
 - [20] Behnood, A. and Mannering, F.L. (2016). "An empirical assessment of the effects of economic recessions on pedestrian-injury crashes using mixed and latent-class models". **Analytic Methods in Accident Research**. 12, 1–17.

- [21] Berman, S. M. (1963). "Note on extreme values, competing risks and semi-Markov processes". **The Annals of Mathematical Statistics**. 34, 1104-1106.
- [22] Bernardinelli, L., Clayton, D. G., Pascutto, C., Montomoli, C., Ghislandi, M., and Songini, M. (1995). "Bayesian analysis of space-time variation in disease risk". **Statistics in Medicine**. 21-22(14), 2433-2443.
- [23] Besag J. (1974). "Spatial interaction and the statistical analysis of lattice systems (with discussion)". **Journal of Royal Statistical Society**. Ser B. 36, 192-236.
- [24] Bhatt, S., Weiss, D. J., Cameron, E., Bisanzio, D., Mappin, B., Dalrymple, U., Battle, K. E., Moyes, C. L., Henry, A., Eckho, P. A., Wenger, E. A., Briet, O., Penny, M. A., Smith, T. A., Bennett, A., Yukich, J., Eisele, T. P., Grin, J. T., Fergus, C. A., Lynch, M., Lindgren, F., Cohen, J. M., Murray, C. L. J., Smith, D. L., Hay, S. I., Cibulskis, R. E., and Gething, P. W. (2015). "The effect of malaria control on plasmodium falciparum in Africa between 2000 and 2015". **Nature**. (526), 207-211.
- [25] Bivand, R., Hauke, J., and Kossowski, T. (2013a). "Computing the Jacobian in Gaussian spatial autoregressive models: an illustrated comparison of available methods". **Geographical Analysis**. 45(2), 150-179.
- [26] Bivand, R. and Lewin-Koh, N. (2017). "maptools: Tools for Reading and Handling Spatial Objects". **R package version 0.9-2**.
- [27] Bivand, R. and Piras, G. (2015). "Comparing implementations of estimation methods for spatial econometrics". **Journal of Statistical Software**. Articles, 63(18).
- [28] Blangiardo, M. and Cameletti, M. (2015). Spatial and Spatio-temporal Bayesian Models with R-INLA. **John Wiley and Sons**.
- [29] Blei, D.M., Kucukelbir, A. and McAuliffe, J.D. (2017). "Variational inference: a review for statisticians". **Journal of the American Statistical Association**. 112(518), 859-877.
- [30] Bolin, D. and Lindgren, F. (2015). "Excursion and contour uncertainty regions for latent Gaussian models". **Journal of the Royal Statistical Society: Series B (Statistical Methodology)**. 77(1), 85-106.

-
- [31] Bolin, D. and Lindgren, F. (2017). "Quantifying the uncertainty of contour maps". **Journal of Computational and Graphical Statistics**. 26(3), 513-524.
- [32] Bosco, C., Alegana, V., Bird, T., Pezzulo, C., Bengtsson, L., Sorichetta, A., Steele, J., Hornby, G., Ruktanonchai, C. and Ruktanonchai, N. (2017). "Exploring the high-resolution mapping of gender-disaggregated development indicators". **Journal of The Royal Society Interface**. 14(129), 2016-0825.
- [8] Boudreau, S. A., Shackell, N. L., Carson, S., and den Heyer, C. E. (2017). "Connectivity, persistence, and loss of high abundance areas of a recovering marine fish population in the north west atlantic ocean". **Ecology and evolution**. 7(22), 9739-9749.
- [33] Box, G. E. P. and Tiao, G. C. (1973). "Bayesian Inference in Statistical Analysis". **Addison-Wesley Publishing Company**. Reading, Mass.-London-Don Mills.
- [34] Box, G., Jenkins, G., M. and Reinsel, G., C. (1994). "Time Series Analysis: Forecasting and Control". third edition. **Prentice-Hall**.
- [35] Braitman, K.A., Kirley, B.B., Ferguson, S. and Chaudhary, N.K. (2007). "Factors leading to older drivers' intersection crashes". **Traffic Injuries Prevention**. 8(3), 267-274.
- [36] Burgette, F.L., Puelz, D. and Hahn, R.P. (2020). "A symmetric prior for multinomial probit models". **Bayesian Analysis**. 1-18.
- [37] Carson, J. and Mannering, F. (2001). "The effect of ice warning signs on ice-accident frequencies and severities". **Accident Analysis and Prevention**. 33(1), 99-109.
- [38] Chen, Z. and Lynn, K. (2001). "Note on the Estimation of the Multinomial Logit Model with Random Effects". **Biometrika**. 55(2), 89-95.
- [39] Cressie, Noela. C., (1993). "Statistics for spatial data". **John wiley and Sons, INC**. New York.
- [40] Clayton, D. G. (1996). "Markov chain Monte Carlo in practice". **Chapman and Hall**.
- [41] Congdon, P. (2005). "Bayesian Models for Categorical Data". **Queen Mary, University of London**. UK.

- [42] Conaway, M. R. (1989). "Analysis of repeated categorical measurements with conditional likelihood methods". **Journal of American Statistic Association**. 84, 53-62.
- [43] Conaway, M. R.(1992). "The analysis of repeated categorical measurements subject to nonignorable nonresponse". **Journal of American Statistic Association**. 87, 817-824.
- [44] Cowles, MK. and Carlin, BP. (1996). "Markov chain Monte Carlo convergence diagnostics: a comparative review". **Journal of American Statistic Association**. 91, 883-904.
- [45] Cox, D. R. (1972). "Regression models and life-tables (with discussion)". **Journal of Royal Statistical Society**. B, 34, 187-220.
- [46] Cressie, N. and Holland, P. (1983). "Characterizing the manifest probabilities of latent trait models". **Psychometrika**. 48, 129-141.
- [47] Cressie, N. and Christopher K. Wikle. (2011). "Statistics for Spatio-Temporal Data". **John willey and son**. New Jercey.
- [48] Cross, Manski, C. (2002). "Regression, short and long". **Econometrica**. 70.
- [49] Czepiel, S. A. (2019). "Maximum likelihood estimation of logistic regression models: Theory and implementation".<https://czep.net/stat>.
- [50] Daneil A.powers, and Y. Xie (1999). "Statistical methods for categorical data analysis". **Academic Press**.
- [51] Darroch, J. N. (1981). "The Mantel-Haenszel test and tests of marginal symmetry, fixed-effects and mixed models for a categorical response". **International Statistics Review**. 49, 285-307.
- [52] Depraetere, N. and Vandebroek, M. (2017). "A comparison of variational approximations for fast inference in mixed logit models". **Computational Statistics**. 32(1), 93-125.
- [53] Dong, C., Richards, S.H., Huang, B. and Jiang, X. (2013). "Identifying the factors contributing to the severity of truck-involved crashes". **International Journal of Injeries Control Safety Promotion**. 22(2), 116-126.

-
- [54] Dow, J. K. and Endersby, J. W. (2004). "Multinomial probit and multinomial logit: a comparison of choice models for voting research". **Electoral Studies**. 23(1), 107-222.
- [55] Eluru, N., Bagheri, M., Miranda-Moreno, L.F. and Fu, L. (2012). "A latent class modeling approach for identifying vehicle driver injury severity factors at highway railway crossings". **Accident Analysis and Prevention**. 47, 119-127.
- [56] Farmer, C.M., Braver, E.R. and Mitter, E.L. (1997). "Two-vehicle side impact crashes: the relationship of vehicle and crash characteristics to injury severity". **Accident Analysis and Prevention**. 29(3), 399-406.
- [57] Fahrmeir L and Tutz G (2001). "Multivariate Statistical Modelling Based on Generalized Linear Models". **Springer-Verlag**. New York.
- [58] Feng, S., Li, Z., Ci, Y. and Zhang, G. (2016). "Risk factors affecting fatal bus accident severity: their impact on different types of bus drivers". **Accident Analysis and Prevention**. 86, 29-39.
- [59] Freni-Sterrantino, A., Ventrucci, M. and Rue, H. (2017). "Note on intrinsic conditional autoregressive models for disconnected graphs". **Spatial and Spatio-temporal Epidemiology**. 26, 25-34.
- [60] Fuglstad, G. A., Simpson, D., Lindgren, F., and Rue, H. (2017). "Constructing priors that penalize the complexity of Gaussian random fields". **Journal of the American Statistical Association**. arXiv:1503.00256.
- [61] Gakidou, E., Afshin, A., Abajobir, A. A., Abate, K. H., Abbafati, C., Abbas, K. M., Abd-Allah, F., Abdulle, A. M., Abera, S. F. and Aboyans, V. (2017). "Global, regional, and national comparative risk assessment of 84 behavioural, environmental and occupational, and metabolic risks or clusters of risks, 1990-2016: a systematic analysis for the global burden of disease study (2016)". **Lancet**. 390(10100), 1345-1422.
- [62] Ghosh M., Zhang L. and Mukherjee B. (2006). "Equivalence of posteriors in the Bayesian analysis of the multinomial Poisson transformation". **Metron-International Journal of Statistics**. 64, 19-28.

- [63] Geedipally S., Turner P. and Patil S. (2011). "Analysis of motorcycle crashes in Texas with multinomial logit model". **Transportation Research Record: Journal of the Transportation Research Board**. 2265, 62–69.
- [64] Golding, N., Burstein, R., Longbottom, J., Browne, A. J., Fullman, N., Osgood-Zimmerman, A., Earl, L., Bhatt, S., Cameron, E. and Casey, D. C. (2017). "Mapping under-5 and neonatal mortality in Africa, 200015: a baseline analysis for the sustainable development goals". **The Lancet**. 390(10108), 2171-2182.
- [65] Goodman, L. A. (1959). "Some alternatives to ecological correlation". **American Journal of Sociology**. 64 (6), 610-625.
- [66] Guse, B., Kiesel, J., Pfannerstill, M. and Fohrer, N. (2020). "Assessing parameter identifiability for multiple performance criteria to constrain model parameters". **Hydrological Sciences Journal**. 2150-3435.
- [67] Gustafson, P. (2004). "On model expansion, model contraction, identifiability and prior information: Two illustrative scenarios involving mismeasured variables". **Statistical Science**. 20(2), 111-140.
- [68] Hall, D. B. and Prstgaard, J. T. (2001). "Order restricted score tests for homogeneity in generalised linear and nonlinear mixed models". **Biometrika**. 88, 739-751.
- [69] Hao, W. and Daniel, J. (2014). "Motor vehicle driver injury severity study under various traffic control at highway-rail grade crossings in the United States". **Journal of Safety Research**. 51, 41–48.
- [70] Hasti, T., J. and Tibshirani, R., J. (1990). Generalized additive models. **Chpman and Hall**.
- [71] Hastings, W.K. (1970). "Monte Carlo Sampling Methods Using Markov Chains and Their Applications". **Biometrika**. 57 (1), 97–109.
- [72] Hausman J and McFadden D. (1984). "Specification tests for the multinomial logit model". **Econometrica**. 52, 1219–1240.
- [73] Heaton, M. J., Datta, A., Finley, A., Furrer, R., Guhaniyogi, R., Gerber, F., Gramacy, R. B., Hammerling, D., Katzfuss, M. and Lindgren, F. (2017). "Methods for analyzing large spatial data: A review and comparison". **arXiv preprint arXiv:1710.05013**.

-
- [74] Held, L. and Rue, H. (2010). "Conditional and intrinsic autoregressions". Handbook of Spatial Statistics. **CRC/Chapman and Hall**. M., and Guttorp, P., editors, Raton, FL.
- [75] Hsiao, C. (1983). "Identification". **Handbook of econometrics** 1. 223-283.
- [76] Hosmer D., Jr., Lemeshow, S. and Sturdivant, R. (2013). "Applied Logistic Regression". **John Wiley Sons Inc**. Hoboken New Jersey.
- [77] Hurwicz, L. (1950). "Generalization of the concept of identification". Statistical Inference in Dynamic Economic Models. **John Wiley and Sons, Inc**.
- [78] Jacquez, J. A. and T. Perry (1990). "Parameter estimation: local identifiability of parameters". **American Journal of Physiology-Endocrinology and Metabolism**. 258 (4), E727-E736.
- [79] Johnson, W. O., Gastwirth, J. L. and Pearson, L. M. (2001). "Screening without a gold standard: The Hui-Walter paradigm revisited". **American Journal of Epidemiology**. 153, 921-924.
- [80] Jain, A., Mukhopadhyay, S., Tolety, S. and Bhattacharjee, T. (2013). "Multinomial logit models, Econometrics II term paper". <http://www.igidr.ac.in/faculty>.
- [81] Jones, S., Gurupackiam, S. and Walsh, J. (2013). "Factors influencing the severity of crashes caused by motorcyclists: analysis of data from Alabama". **Journal of Transportation Engineering, Part B: Pavements**. 139, 949-956.
- [82] Jousimo, J., Tack, A. J. M., Ovaskainen, O., Mononen, T., Susi, H., Tollenaere, C., and Laine, A.-L. (2014). "Ecological and evolutionary effects of fragmentation on infectious disease dynamics". **Science**. 344(6189), 1289-1293.
- [83] Khattak, A. (2001). "Injury severity in multivehicle rear-end crashes". **Transportation Research Record: Journal of the Transportation Research Board**. 1746, 59-68.
- [84] Knorr-Held, L. (2000). "Bayesian modelling of inseparable space-time variation in disease risk". **Statistics in Medicine**. 19(17-18), 2555-2567.
- [85] Kockelman, K. and Kweon, J.Y. (2002). "Driver injury severity: an application of ordered probit models". **Accident Analysis and Prevention**. 34(3), 313-321.

- [86] Koopmans, T. C. (1949). "Identification problems in economic model construction". **Econometrica, Journal of the Econometric Society**. 125-144.
- [87] Koopmans, T. C. and O. Reiersol (1950). "The identification of structural characteristics". **The Annals of Mathematical Statistics**. 21 (2), 165-181.
- [88] Krainski, E.T. (2018). **Statistical Analysis of Space-time Data: New Models and Applications**. PhD thesis, Norwegian University of Science and Technology.
- [89] Kudo, A. (1963). "A multivariate analogue of the one-sided test". **Biometrika**. 50, 403-418.
- [90] Kweon, Y.J. and Kockelman, K.M. (2003). "Overall injury risk to different drivers: combining exposure, frequency, and severity models". **Accident Analysis and Prevention**. 35(4), 441-450.
- [91] Laplace, P S. (1774). "Mémoires de mathématique et de physique, Tome Sixième" Memoir on the probability of causes of events". **Statistical Science**. 1(3), 366-367.
- [92] Lee, J. and Mannering, F. (2002). "Impact of roadside features on the frequency and severity of run-off-roadway accidents: an empirical analysis". **Accident Analysis Prevention**. 34, 149-161.
- [93] Lee J., Yasmin S., Eluru N., Abdel-Aty M. and Cai, Q. (2018). "Analysis of crash proportion by vehicle type at traffic analysis zone level: a mixed fractional split multinomial logit modeling approach with spatial effects". **Accident Analysis Prevention**. 111, 12-22.
- [94] Lee L.Y.J., Green J.P. and Ryan M.L. (2017). "On the "Poisson Trick" and its extensions for fitting multinomial regression models". **arXivpreprint arXiv:1706.09523**.
- [95] Lewis, B. (2003). "Small Dams". CRC Press/Balkema.
- [96] Lin. X. (1997). "Variance component testing in generalised linear models with random effects". **Biometrika**. 84, 309-326.
- [97] Liu, J., Khattak, J.A., Li, X., Nie, Q. and Ling, Z. (2020). "Bicyclist injury severity in traffic crashes: a spatial approach for geo-referenced crash data to uncover non-stationary correlates". **Journal of Safety Research**. 73, 25-35.

-
- [98] Luce, R. D. (1959). "Individual choice behaviour". **John Wiley**. New York.
 - [99] Maddala, G.S., (1983). "Limited-Dependent and Qualitative Variables in Econometrics". **Cambridge University Press**. New York.
 - [100] Martin, AD., Quinn. KM. and Park, JH. (2011). "MCMCpack: Markov Chain Monte Carlo in R". **Journal of Statistical Software**. 42(9), 22.
 - [101] Matzkin, R. L. (2007). "Nonparametric identification". **Handbook of Econometrics**. 6, 5307-5368.
 - [102] Manski, C. F. (2009). Identification for prediction and decision. **Harvard University Press**.
 - [103] Martinez-Beneito, M. A., López-Quilez, A., and Botella-Rocamora, P. (2008). "An autoregressive approach to spatio-temporal disease mapping". **Statistics in Medicine**. 27(10), 2874-2889.
 - [104] McCullagh, P. (1982). "Some applications of quasisymmetry". **Biometrika**. 69, 303-308.
 - [105] McCullagh, P., Nelder, J. (1989). "Generalized Linear Models". Second edition. **Chapman and Hall**.
 - [106] McFadden, D. (1981). "Econometric models of probabilistic choice". **Structure Analysis of Discrete Data with Econometric Applications**. MIT Press, Cambridge MA.
 - [107] Mervyn. J. S. (1994). "On tests against one-sided hypotheses in some generalized linear models". **Biometrics**. 50, 853-858.
 - [108] Miaou, S.P. (1994). "Relationship between truck accidents and geometric design of road sections: Poisson vs. negative binomial regressions". **Accident Analysis and prevention**. 26, 471-482.
 - [109] Minasny, B., Vrugt, J.A. and McBratney, A.B. (2011). "Confronting uncertainty in model-based geostatistics using Markov Chain Monte Carlo simulation". **Geoderma**. 163, 150-162.
 - [110] Milledge, D.G., Lane, S.N., Heathwaite, A.L. and Reaney, S.M. (2012). "A Monte Carlo approach to the inverse problem of diffuse pollution risk in agricultural catchments". **Science of the Total Environment**. 433, 434-449.

- [111] Mohammed, A.A., Ambak, K., Mosa, A.M. and Syamsunur, D. (2019). "A review of the traffic accidents and related practices worldwide". **Open Transportation Journal**. 13, 65-83.
- [112] Moraga, P., Cramb, S. M., Mengersen, K. L., and Pagano, M. (2017). "A geostatistical model for combined analysis of point-level and area-level data using INLA and SPDE". **Spatial Statistics**. 21, 27-41.
- [113] Mullahy, J. and Robert, S. A. (2010). "No time to lose: Time constraints and physical activity in the production of health". **Review of the Economics of the Household**. 8, 409-432.
- [114] Noor, A. M., Kinyoki, D. K., Mundia, C. W., Kabaria, C. W., Mutua, J. W., Alegana, V. A., Fall, I. S., and Snow, R. W. (2014). "The changing risk of Plasmodium falciparum malaria infection in Africa: 2000-10: a spatial and temporal analysis of transmission intensity". **Lancet**. 383(9930), 1739-1747.
- [115] O'Donnell, C.J. and Connor, D.H. (1996). "Predicting the severity using models of ordered multiple choice". **Accident Analysis and Prevention**. 28(6), 739 - 753.
- [116] Ormerod, J.T. and Wand, M.P. (2010). "Explaining variational approximations". **Journal of the American Statistical Association**. 64(2), 140-153.
- [117] Pahlavan-Rad, M.R.; Khormali, F.; Toomanian, N.; Brungard, C.W.; Kiani, F.; Komaki, C.B. and Bogaert, P. (2016). "Legacy soil maps as a covariate in digital soil mapping: A case study from Northern Iran". **Geoderma**. 279, 141-148.
- [118] Palmgren, J. (1981). "The Fisher information matrix for log linear models arguing conditionally on observed explanatory variables". **Biometrika**. 68, 563-566.
- [119] Palmgren, J. and Ekholm, A. (1987). "Exponential family non-linear models for categorical data with errors of observation". **Applied Stochastic Models and Data Analysis**. 3, 111-124.
- [120] Paulino, C. D. M. and C. A. de Braganca Pereira (1994). "On identifiability of parametric statistical models". **Journal of the Italian Statistical Society**. 3(1), 125-151.
- [121] Penmetsa, P. and Pulugurtha, S.S. (2017). "Modeling crash injury severity by road feature to improve safety". **Traffic Injury Prevention**. 19, 102-109.

- [122] Pettit, L.I. (1990). "The conditional predictive ordinate for the normal distribution". **Journal of the Royal Statistical Society. Series B (Methodological)**. 52(1), 175-84.
- [123] Pirdashti, H., Pirdashti, M., Mohammadi, M., Gharavi Baigi, M. and Movagharnejad, K. (2015). "Efficient use of energy through organic rice-duck mutualism system". **Agronomy for Sustainable Development**. 35, 1489-1497.
- [124] Plummer, M. (2016). "rjags: Bayesian Graphical Models using MCMC". **R package version 4-6**.
- [125] Porter, B.E. (2011). "Handbook of Traffic Psychology". **Old Dominion University Norfolk VA. USA, Elsevier**.
- [126] Quinn, K.M., Martin, A.D. and Whitford, A.B., (1999). "Voter choice in multi-party democracies: a test of competing theories and models". **American Journal of Political Science**. 43, 1231-1247.
- [127] R Core Team. (2021). "R: A language and environment for statistical computing, R Foundation for Statistical Computing". **Vienna, Austria. URL <https://www.R-project.org/>**.
- [128] Regev, S., Rolison, J.J. and Moutari, S. (2018). "Crash risk by driver age, gender, and time of day using a new exposure methodology". **Journal of Safety Research**. 66, 131-140.
- [129] Rezapour, M., Molan, M.A. and Ksaibati, K. (2019). "Application of multinomial regression model to identify parameters impacting traffic barrier crash severity". **The Open Transportation Journal**. 13, 57-64.
- [130] Robert, C. P. and Casella, G. (1999). "Monte Carlo Statistical Methods". **Springer-Verlag**. New York.
- [131] Rothenberg, T. J. (1971). "Identification in parametric models". **Econometrica: Journal of the Econometric Society**. 577-591.
- [132] Rue, H., Held, L. (2005). "Gaussian Markov Random Fields: Theory and Applications, Monographs on Statistics and Applied Probability". **Chapman and Hall**. Boca Raton.

- [133] Rue, H., Martino, S., Chopin, N. (2009). "Approximate Bayesian inference for latent Gaussian models by using integrated nested Laplace approximations". **Journal of Royal Statistical Society. Series B (Statistical Methodology)**. 71, 319-392.
- [134] Rue, H., Riebler, A., Sigrunn, H., Sørbye, J., Illian, 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.
- [135] Ruiz-Cárdenas, R., Krainski, E. T. and Rue, H. (2012). "Direct fitting of dynamic models using integrated nested Laplace approximations-INLA". **Computational Statistics and Data Analysis**. 56(6), 1808-1828.
- [136] Rushworth, A., Lee, D. and Mitchell, R. (2014). "A spatio-temporal model for estimating the long- term effects of air pollution on respiratory hospital admissions in Greater London". **Spatial and Spatio-temporal Epidemiology**. 1(10), 29-38.
- [137] Rom., M., t. and Cohen, A., (1995). "Estimation in the polytomous logistic regression model". **Journal of Statistical Planning and Inference**. 43, 341-353.
- [138] Sadeghi-Bazargani, H., Sharifian, S., Khorasani-Zavareh, D., Zakeri, R., Sadigh, M., Golestani, M., Masoudifar, R., Rahmani, F., Mikaeeli, N., Namvaran, J., Pour-Ebrahim, K., Rezaei, M., Arabzadeh, B., Samadirad, B., Seyffarshad, A., Mirza-Mohammadi-Teimorloue, F., Kazemnezhad, S., Marin, S., Sheikhi, S. and Mohammadi, R. (2020). "Road safety data collection systems in Iran: a comparison based on relevant organizations". **Chinese Journal of Traumatology**. 23, 265-270.
- [139] Said, F. (2020). **Constrained Statistical Inference for Categorical Data**. PhD thesis. Carleton University, Canada.
- [140] Salum, J.H., Kitali, A.E., Bwire, H., Sando, T. and Alluri, P. (2019). "Severity of motorcycle crashes in Dar es Salaam, Tanzania". **Traffic Injuries Prevention**. 20(2), 189-195.
- [141] Savolainen, P., Mannering, F. (2007). "Probabilistic models of motorcyclists' injury severities in single- and multi-vehicle crashes". **Accident Analysis and Prevention**. 39(5), 955-963.
- [142] Savolainen, P.T., Mannering, F.L., Lord, D. and Quddus, M.A. (2011). "The statistical analysis of highway crash-injury severities: a review and assessment of methodological alternatives". **Accident Analysis and Prevention**. 43(5), 1666-1676.

-
- [143] Schott., J., R. (2016). "Matrix Analysis for statistics". third edition. **Wiley and Sons, Inc.** Hoboken, New Jersey.
- [144] Schrodle, B., Held, L., and Rue, H. (2012). "Assessing the impact of a movement network on the spatio-temporal spread of infectious diseases". **Biometrics**. 68(3), 736-744.
- [145] Seaman, S. R. and Richardson, S. (2004). "Equivalence of prospective and retrospective models in the Bayesian analysis of case-control studies". **Biometrika**. 91, 15-25.
- [146] Serafini, F. (2018). "Multinomial logit models with INLA: vignettes for INLA". <https://rdrr.io/github/inbo/INLA/f/vignettes/multinomial.Rmd>.
- [147] Shaddick, G., Thomas, M. L., Green, A., Brauer, M., Donkelaar, A., Burnett, R., Chang, H. H., Cohen, A., Dingenen, R. V. and Dora, C. (2018). "Data integration model for air quality: a hierarchical approach to the global estimation of exposures to ambient air pollution". **Journal of the Royal Statistical Society**, Series C (Applied Statistics), 67(1), 231-253.
- [148] Shaheed, M., Gkritza, K., Zhang, W. and Hans, Z. (2013). "A mixed logit analysis of two-vehicle crash severities involving a motorcycle". **Accident Analysis and Prevention**. 61, 119-28.
- [149] Simpson, D. P., Rue, H., Riebler, A., Martins, T. G., and Sørbye, S. H. (2017). "Penalising model component complexity: A principled, practical approach to constructing priors (with discussion)". **Statistical Science**. 32(1), 1-28.
- [150] Slessarev E. W., Lin Y., Bingham N. L., Johnson J. E., Dai Y. and Schimel J. P. (2016). "Water balance creates a threshold in soil pH at the global scale". **Nature**. 540(7634): 567.
- [151] Sørbye, S. and Rue, H. (2014). "Scaling intrinsic Gaussian Markov random field priors in spatial modelling". **Spatial Statistics**. 8, 39-51.
- [152] Soroori, E., Moghaddam, M., A. and Salehi, M. (2020). "Modelling spatial non-stationary and overdispersed crash data: development and comparative analysis of global and geographically weighted regression models applied to macrolevel injury crash data". **Journal of Transportation Safety and Security**. 13, 1000-1024.

- [153] Spiegelhalter, D. J., Thomas, A., Best, N. G., and Gilks, W. R. (1995). "BUGS: Bayesian inference using Gibbs sampling". **Version 0.50, MRC Biostatistics Unit, Cambridge.**
- [154] Stan Development Team (2015). "Stan Modelling Language User's Guide and Reference Manual". **Version 2.6.1.**
- [155] Sun, Y., Wang, Y., Yuan, K., Chan, O., T. and Huang, Y. (2020). "Discovering spatio-temporal clusters of road collisions using the method of fast Bayesian model-based cluster detection". **Sustainability**. 12(20), 81-86.
- [156] Swartz, T., Haitovsky, Y., vexler., A., yang, T. (2004). "Bayesian identifiability and misclassification in multinomial data". **The Canadian Journal of Statistics**. 32, 1-18.
- [157] Theil, H. (1969). "A multinomial extension of the linear logit model". **International Economic Review**. 10, 251-259.
- [158] Ugarte, M. D., Adin, A., Goicoa, T., and Militino, A. F. (2014). "On fitting spatio-temporal disease mapping models using approximate Bayesian inference". **Statistical methods in medical research**. 23(6), 507-530.
- [159] Vivar, J. C. and Ferreira, M. A. R. (2009). "Spatio-temporal models for Gaussian areal data". **Journal of Computational and Graphical Statistics**. 18(3), 658-674.
- [160] Wakefield, J., Fuglstad, G.-A., Riebler, A., Godwin, J., Wilson, K., and Clark, S. J. (2018). "Estimating under five mortality in space and time in a developing world context". **Statistical Methods in Medical Research**. 28(9), 2614-2634.
- [161] Washington, S., Karlaftis, M. and Mannering, F. (2003). "Statistical and Econometric Methods for Transportation Data Analysis". **Chapman and Hall/CRC.**
- [162] Wedagama, D. and Dissanayake, D. (2020). "A model of latent class multinomial logit to investigate motorcycle accident injuries". **Engineering and Applied Science Research**. 47(4), 422-429.
- [163] Whitehead, J. (1980). "Fitting Cox's regression model to survival data using GLIM". **Journal of Applied Statistics**. 29, 268-275.

-
- [164] Wieland, F., Hauber, L., A., Rosenblatt, M., Tönsing, C. and Timmer, J. (2021). "On structural and practical identifiability". **Current Opinion in Systems Biology**. 25, 60-69.
- [165] World Health Organization (2016). **Ambient air pollution: A global assessment of exposure and burden of disease**.
- [166] Xie, Y., Zhang, Y. and Liang, F. (2009). "Crash injury severity analysis using Bayesian ordered probit models". **Journal of Transportation Engineering**. 135(1), 18-25.
- [167] X. Lin. (1997). "Variance component testing in generalised linear models with random effects". **Biometrika**. 84, 309-326.
- [168] Yamamoto, T. and Shankar, V. (2004). "Bivariate ordered response probit model of driver's and passenger's injury severities in collisions with fixed object". **Accident Analysis Prevention**. 36(5), 869-876.
- [169] Zeng, Q., Wen, H. and Huang, H. (2016). "The interactive effect on injury severity of driver vehicle units in two-vehicle crashes". **Journal of Safety Research**. 59, 105-111.
- [170] Zeng, Q., Gu, W., Zhang, X., Wen, H., Lee, J. and Hao, W. (2019). "Analyzing freeway crash severity using a Bayesian spatial generalized ordered logit model with conditional autoregressive priors". **National Library of Medicine**. 127, 87-95.
- [171] Zhang, H. (2004). "Inconsistent estimation and asymptotically equal interpolations in model-based geostatistics". **Journal of the American Statistical Association**. 99(465), 250-261.
- [172] Zhang Y. (2011). "Exploring single vehicle crash severity on Rural, two-lane highways with crash-level and occupant-level multinomial logit models, doctoral dissertation". **Salt Lake City, UT: Department of Civil and Environmental Engineering**. University of Utah.
- [173] Ziakopoulos, A. and Yannis, G. (2020). "A review of spatial approaches in road safety". **Accident Analysis and Prevention**. 135, 105-323.

واژه‌نامه فارسی به انگلیسی

Hyper parameter	ابریارامتر
Effect	اثر
Latent effect	اثر پنهان
Fixe effect	اثر ثابت
Spatail effect	اثر فضایی
Probability	احتمال
Cumulative probability	احتمال تجمعی
Inferance	استنباط
Nominal	اسمی
Bayesian Inferance	استنباط بیزی
Trail	آزمایش
Nested	بیز
Estimate	برآورد
Estimator	برآوردگر
Random vector	برداری تصادفی
Odds	بخت‌ها
High dimentional	بعد بالا
Bayes	بیز
Associated prime ideal	بیز تجربی
Parameter	پارامتر
Continues	پیوسته
Predictor	پیشگو
Probit	پروبیت
Prior	پیشین
Poisson	پوآسن

Link function	تابع پیوند
Likelihood function	تابع درست‌نمایی
Transformation	تبدیل
Cumulative	تجمعی
Approximation	تقریب
Laplace Approximation	تقریب
Random	تصادفی
Ordinal	ترتیبی
Analysis	تحلیل
Distribution	توزیع
Beta Distribution	توزیع بتا
Prior distribution	توزیع پیشین
Posterior distribution	توزیع پسین
Marginal posterior distribution	توزیع پسین حاشیه‌ای
Stationary distribution	توزیع مانا
Contingency table	جدول احتمالی
Density	چگالی
Multinomial	چندجمله‌ای
Multi-category	چند رده‌ای
Multi-variable	چند متغیره
Exponential family	خانواده نمایی
Linear	خطی
Lattice data	داده شبکه‌ای
Likelihood	درست‌نمایی
Binomial	دوجمله‌ای
Binary	دودویی
Categorical	رده‌ای
Longitudinal	طولی
Intercept	عرض از مبدا
Non-linear	غیر خطی
Gamma	گاما
Spatial	فضایی

Spatio-temporal	فضایی-زمانی
Closed form	فرم بسته
Covariance	کوواریانس
Cumulative logit	لجیت تجمعی
Log-likelihood	لگاریتم درست‌نمایی
Log-linear	لگ-خطی
Markovian	مارکوفی
Maximum	ماکزیمم
Maximum likelihood	ماکزیمم درست‌نمایی
Linear probability model	مدل احتمال خطی
Metropolis-Hasting	متروپولیس-هاستینگ
variable	متغیر
Response variable	متغیر پاسخ
Latent variable	متغیر پنهان
Random variable	متغیر تصادفی
explanatory variable	متغیر توضیحی
Parametric model	مدل پارامتری
Non-Parametric model	مدل ناپارامتری
Generalized linear model	مدل خطی تعمیم یافته
Hierarchical model	مدل سلسله مراتبی
Independence	مستقل
Real value	مقدار حقیقی
Gaussian markov random field	میدان تصادفی گوسی مارکوفی
Random field	میدان تصادفی
Normal	نرمال
Odds ratio	نسبت بخت
Sample	نمونه
sampling	نمونه‌گیری

واژه‌نامه انگلیسی به فارسی

Analysis	تحلیل
Approximation	تقریب
Bayes	بیز
Bayesian inference	استنباط بیزی
Beta distribution	توزیع بتا
Binary	دودویی
Binomial	دوجمله‌ای
Categorical	رده‌ای
Contingency table	جدول احتمالی
Continues	پیوسته
Covariance	کوواریانس
Cononical link	پیوند کانونی
Closed form	فرم بسته
Cumulative	تجمعی
Cumulative logit	لجیت تجمعی
Cumulative probability	احتمال تجمعی
Density	چگالی
Distribution	توزیع
Empirical bayes	بیز تجربی
Estimate	برآورد
Estimator	برآوردگر
Explanatory variable	متغیر توضیحی
Function	تابع
Gamma	گاما
Gaussian markov random field	میدان تصادفی مارکوفی گوسی

Genralized linear model	مدل خطی تعمیم‌یافته
Gibbs sampling	نمونه‌گیری گیبس
High dimetional	بعد بالا
Hierarical model	مدل سلسله مراتبی
Hyper parameter	ابر پارامتر
Independance observations	مشاهدات مستقل
Intercept	عرض از مبدا
Latent effect	اثر پنهان
Lattice data	داده‌های شبکه‌ای
Log-likelihood	لگاریتم درست‌نمایی
Log-linear	لگ-خطی
Lognitudinal	طولی
Link function	تابع پیوند
Miancearginal response	پاسخ حاشیه‌ای
Marginal model	مدل حاشیه‌ای
Multinomial	چندجمله‌ای
Parametric model	مدل پارامتری
Random variable	متغیر تصادفی
Prior	پیشین
Posterior	پسین
Sample	نمونه
Sampling	مونه‌گیرین

Abstract

A categorical variable is a random variable that groups observations. For example, classifying people with different characteristics according to their blood type results in a categorical variable. Nowadays, the widespread of various sciences has made analyzing categorical data very essential. Hence, different scientific disciplines pursue more accurate methods for analyzing such data. In classical statistics, several methods have been proposed to analyze categorical data. These existing methods and relevant software packages are not efficient enough for analyzing data with complex spatial and spatio-temporal structures. Bayesian analysis and inference of categorical data require sampling methods such as the MCMC algorithms, while they are not also efficient for categorical data with complex structures. Our proposed approach, in this thesis, is to use an approximate Bayesian method, called integrated nested Laplace approximation (INLA), which is not accompanied by the significant problems seen in the MCMC algorithms for implementing the Bayesian analysis of spatial models. To deal with the inherent identifiability problem of the multinomial models, we consider two different types of identifiability constraints and compare their performances. We use the Poisson-multinomial transformation to develop the model in the class of latent Gaussian models applicable for INLA. We also utilize the individualized logit model as another type of modeling depending on the response variable type. Finally, we assess and compare the proposed models through both simulated and real data examples.

Keywords: Bayesian inference; fractional split multinomial model; identifiability; individualized logit; INLA method; multinomial-Poisson transformation; spatial multinomial model; spatio-temporal dependency structure.



Faculty of Mathematical Sciences

PhD Thesis in Statistics

Bayesian Categorical Data Analysis with Spatial and Spatio-temporal Structures

By: Leila Barmoudeh

Supervisor:
Dr. Hossein Baghishani

February 2022