

تحليل داده‌های عصبی

مؤلفان:

رابرت ای. کاس

اُوری تی. اِدِن

امری ان. براون

مترجمان:

سیدجمال میرکمالی

فرزانه صفوی‌منش



انتشارات دانشگاه اراک

سرشناسه :	Kass, Robert E - کاس، رابرت
عنوان و نام پدیدآور :	تحلیل داده‌های عصبی / مولفان رابرت ای. کاس، اوری تی. ایدن، امری ان. براون؛ مترجمان سیدجمال میرکمالی، فرزانه صفوی‌منش
مشخصات نشر :	اراک: دانشگاه اراک، انتشارات، ۱۴۰۱.
مشخصات ظاهری :	۱۸۱ ص.: مصور، نمودار.
شابک :	۹۷۸-۶۲۲-۹۲۷۴۵-۴-۵
وضعیت فهرست‌نویسی :	فیپا
یادداشت :	عنوان اصلی: Analysis of neural data, 2014.
موضوع :	تحلیل گره‌های عصبی Neural analyzers روان‌شناسی عصبی - آزمون‌ها Neuropsychological tests عصب پایه‌شناسی - ریاضیات Neurosciences - Mathematics
شناسه افزوده :	ایدن، یوری تی.
شناسه افزوده :	Eden, Uri T.
شناسه افزوده :	براون، امری ان.
شناسه افزوده :	Brown, E. N. (Emery N.)
شناسه افزوده :	میرکمالی، سیدجمال، ۱۳۶۳-، مترجم
شناسه افزوده :	صفوی‌منش، فرزانه، ۱۳۶۳-، مترجم
شناسه افزوده :	دانشگاه اراک. انتشارات
رده‌بندی کنگره :	QP۳۵۷/۵
رده‌بندی دیویی :	۶۱۲/۸
شماره کتابشناسی ملی :	۹۰۹۹۸۲۷

این کتاب مشمول قانون حمایت از حقوق مؤلفان و مصنفان است. تکثیر کتاب به هر روش اعم از فتوکپی، ریسوگرافی، تهیه فایل‌های لوح فشرده، بازنویسی در وبلاگ‌ها، سایت‌ها، مجله‌ها و کتاب، بدون اجازه کتبی ناشر مجاز نیست و موجب پیگرد قانونی می‌شود و تمامی حقوق برای ناشر محفوظ است.

این کتاب با حمایت ستاد توسعه علم و فناوری های شناختی تهیه و چاپ شده است.

عنوان: تحلیل داده‌های عصبی

تألیف: رابرت ای. کاس، اوری تی. ایدن، امری ان. براون

ترجمه: سیدجمال میرکمالی، فرزانه صفوی‌منش

نوبت چاپ: اول

تاریخ انتشار: ۱۴۰۲

شماره گان: ۲۰۰ نسخه

ناشر: انتشارات دانشگاه اراک

چاپ و صحافی: انتشارات دانشگاه اراک

«مسئولیت صحت مطالب کتاب با مؤلفان است»

قیمت: ۶۰۰ هزار تومان

اراک، میدان بسیج، بلوار کربلا، دانشگاه اراک، ساختمان کتابخانه مرکزی و مرکز اسناد، طبقه دوم، اتاق شماره

۲، انتشارات دانشگاه اراک پست الکترونیک: press@araku.ac.ir - تارنما: https://press.araku.ac.ir



پیشگفتار مترجمان

وقتی جناب آقای دکتر کمال خرازی، رئیس وقت ستاد علوم شناختی، پیشنهاد ترجمه این کتاب را مطرح کردند، ویژگی‌های منحصر به فرد کتاب، در مقایسه با انگشت‌شمار کتاب‌های متفرقه منتشر شده در حوزه‌های مرتبط با مباحث این کتاب، مترجمان را بر آن داشت که بلافاصله پیشنهاد ایشان را بپذیرند و دست‌به‌کار ترجمه کتاب شوند. نبود کتابی جامع برای تدریس درس‌های آمار و روش تحقیق دانشجویان مقاطع دکتری و کارشناسی ارشد گرایش‌های متعدد علوم شناختی از جمله مهم‌ترین دلایلی بود که ترجمه چنین کتابی را ضروری می‌ساخت. همان‌طور که نویسندگان در پیشگفتار کتاب نیز ذکر کرده‌اند، این کتاب حاصل بیش از بیست سال تدریس و پیاده‌سازی روش‌های آماری در تحلیل داده‌های عصبی است. تنوع مثال‌ها به گونه‌ای است که نه تنها پژوهشگران تمامی گرایش‌های علوم شناختی می‌توانند از این کتاب بهره‌مند شوند، بلکه کتاب را به منبعی عالی برای پژوهشگران و دانشجویانی تبدیل کرده است که می‌خواهند با نحوه کاربرد و پیاده‌سازی روش‌های آماری، از سطح مقدماتی تا پیشرفته، آشنا شوند. از آن‌جا که کدهای نرم‌افزاری به کار رفته برای تحلیل داده‌ها به دو زبان R و MATLAB در سایت شخصی نویسندگان به صورت آزاد در اختیار همگان قرار گرفته است، پژوهشگران به راحتی می‌توانند از این کتاب برای آموزش خودخوان مفاهیم آماری و نحوه به کارگیری آن‌ها نیز استفاده کنند.

در انتخاب معادل‌ها برای واژه‌ها و اصطلاحات عمومی، معیار را واژه‌های مصوب کارگروه واژه‌گزینی تخصصی آمار در فرهنگستان زبان و ادب فارسی قرار دادیم. برای انتخاب معادل‌های فارسی واژه‌های تخصصی نیز از واژه‌نامه‌های تخصصی پژوهشکده آمار، واژه‌نامه تخصصی علوم شناختی فرهنگ معاصر و برخی دیگر استفاده کرده‌ایم. هر جا که هنوز معادلی برای واژه‌ای به تصویب نرسیده بود، یا به گواهی تعدد معادل‌ها در کتاب‌های حوزه‌های مرتبط با علوم شناختی اتفاق نظر وجود نداشت، تا جایی که ممکن بود، از دانش و لطف برخی از استادان بخش واژه‌گزینی تخصصی فرهنگستان زبان و نیز برخی سرپرستان مجرب مرکز نشر دانشگاهی کمک گرفتیم. هر جا که به این نتیجه رسیدیم که معادل‌های مصوب هنوز چندان رایج نشده‌اند و ممکن است استفاده از آن‌ها منجر به سردرگمی و مانع از روانی خواندن متن شود، از معادل‌های رایج اما غیرمصوب مصطلح در بین جامعه آماری یا علوم اعصاب استفاده کردیم. در معدود مواردی هم که هیچ کدام از این‌ها را مناسب این کتاب نیافتیم، از صورت فارسی‌نویسی شده آن واژه استفاده کردیم؛ از

آن جمله‌اند: کوواریانس، کورتکس (به‌جای قشر)، فیلتر (به‌جای پالاینده)، بوت‌استرپ (به‌جای خودگردان)، وُکسل، آلل، بافر، یچِ کَلَمپ، آندارترکتومی، رفرکتوری و برخی دیگر. در هر صورت، سعی کردیم که، برای هرچه روان‌تر شدنِ مطالعهٔ متن، معادل‌های انگلیسی واژه‌ها را، گاه همراه با برخی توضیحاتِ بیشتر برای درکِ بهترِ آن‌ها، در پانویس بیاوریم.

گرچه مترجمان طی دو سالی که ترجمهٔ این کتاب به طول انجامید سعی داشتند که در حد وسع و زمان و دسترسی‌های خود جوانب دقت و ظرافت را رعایت کنند، به قول زنده‌یاد کریم امامی در فصلی از کتاب «از پست و بلند ترجمه» اش، «ترجمه راهی‌ست باریک و دشوار که برای گذشتن از آن آدم باید چابکی، نرمش، و حسِ توازنِ یک بندباز را داشته باشد؛ اما چه کسی می‌تواند بی‌آن‌که خودش را به خطر بیفکند از این راه باریک و دشوار به سلامت بگذرد؟». در این مسیر پرخطر، احتمالی آن بسیار است که، به دلیل حجمِ سنگینِ کتاب، یک‌دست‌سازیِ کاملِ برخی از مواردی که برآمده از سبک و سلیقه و انتخابِ متفاوتِ دو مترجم بوده‌اند میسر نشده باشد و مواردی از قلم افتاده باشند. همچنین، ممکن است پیشینهٔ آماری هر دو مترجم برای درک کاملِ مباحثِ علوم اعصاب کافی نبوده و الزاماً به برداشت‌های درستی منتهی نشده باشد. بنابراین، مترجمان از هر پیشنهادی که منجر به کاهشِ ابهاماتِ خوانندگانِ کتاب شود استقبال می‌کنند. تمامیِ اصلاحاتی که پس از سپردنِ کتاب به ناشر محترم لازم باشد، در صفحهٔ اینترنتی اختصاص‌یافته به کتاب توسط مترجمان به نشانی <https://github.com/mirkamali/AND> منتشر خواهد شد. مترجمان کتاب نیز از طریق آدرس‌های ایمیل mirkamali.sj@gmail.com و f.safavimanesh@gmail.com در دسترس و در حدِ وسعِ خود پاسخگوی پیشنهادات شما خواهند بود.

در پایان، از مدیر محترم انتشارات دانشگاه اراک که چاپ کتاب را پذیرفتند، داورانی که با نظرات خود کتاب را خواندنی‌تر ساختند، سرکار خانم کوثر سلیمی که بازخوانی متن و اصلاح حروف‌چینی را بر عهده داشتند، خانم‌ها شقایق بهرامی، محدثه دارابی، و سارا مومنیان که متن‌ها را بازبینی و ویرایش ادبی کردند و نیز دانشجویان و همکارانی که انگیزهٔ ترجمهٔ کتاب را بیشتر کردند، سپاسگزاریم. امیدواریم که کتاب مورد استفادهٔ پژوهشگران حوزه‌های مرتبط با مباحث موجود در آن واقع شود.

فرزانه صفوی‌منش

دکتری آمار از دانشگاه آلبورگ، دانمارک
آماردان ارشد سازمان تحقیقات بالینی Larix،
گروه تحقیقاتی ALTEN، کپنهاگ، دانمارک،
عضو سابق هیئت علمی دانشگاه شهیدبهشتی

سیدجمال میرکمالی

دکتری آمار از دانشگاه شهیدبهشتی
عضو هیئت علمی دانشگاه اراک
دانشکدهٔ علوم پایه،
گروه ریاضی و آمار

پیشگفتار نویسندگان

این کتاب به‌منزله یک راهنما و مرجعی برای هر کسی به‌کار می‌رود که می‌خواهد تحلیل داده‌های عصبی تولیدشده در مطالعاتی را بفهمد که از مولکول‌ها تا مدارها و سیستم‌ها و رفتار گسترده شده‌اند.

ریشه‌های شکل‌گیری آن را می‌توان در تصمیم دو نفر از ما (ای. ان. بی و آر. ای. کی) در ۱۹۹۸ مبنی بر نوشتن یک مقاله مروری در تحلیل آماری داده‌های زنجیر اسپایکی پیدا کرد. اندکی پس از آغاز، تشخیص دادیم که برخی از روش‌هایی که فکر می‌کردیم باید حتماً مرور شده باشد، در واقع، هنوز گسترش نیافته‌اند. پس از آن‌که ما و دیگران این وضع را اصلاح کردیم، دو مقاله مروری منتشر کردیم (براون و همکاران، ۲۰۰۴؛ کاس و همکاران، ۲۰۰۵). در طول این مدت، ما علایق خود را به دیگر مدالیته‌های آزمایشی همچون تصویربرداری عصبی را نیز گسترش دادیم و تدریس در کارگاه‌ها و درس‌های ترمی طولانی‌مدتی را در روش‌های آماری برای علوم اعصاب آغاز کردیم. علاوه بر آن، ما با نویسنده سوم این کتاب (یو، ای) آشنا شدیم که علایق مشترکی با ما در پژوهش و تدریس داشت (و کسی بود که رساله دکتری‌اش را تحت راهنمایی ای. ان. بی انجام داده بود).

بر ما روشن شد که کتابی درباره این موضوع به‌شدت لازم بود و ما توافق کردیم که یکی بنویسم. در حالی که این کار به پروژه‌ای وسیع‌تر از آنچه پیش‌بینی کرده بودیم تبدیل شده بود، همکاری‌های پژوهشی بی‌شمار، گفت‌وگوهایی با همکاران در جلسات، و پیشنهادات وسیعی از دانشجویان بینشی درباره محتوا و ارائه اصول و تکنیک‌هایی بخشید که منجر به این جلد شدند. ما احساس می‌کنیم که ما بسیار آگاه‌تر از زمانی هستیم که شروع کردیم، و امیدواریم که در رساندن بخش زیادی از آنچه در این فرایند یاد گرفته‌ایم موفق بوده باشیم.

برخی خوانندگان ممکن است انتظار کتابی را داشته باشند که بر حسب انواع داده‌های عصبی مرتب شده باشد. ما تصمیم گرفتیم که به جای آن کتاب را بر حسب تحلیل‌ها سازماندهی کنیم، به‌طوری که هر فصل به گستره‌ای از مفاهیم آماری طبقه‌بندی شده اختصاص یابد که مختصراً در عنوان بخش‌ها، که نسخه مبسوط‌تری از فهرست مطالب‌اند، توصیف شده‌اند. با این حال، هر فصل شامل چندین مثال از نحوه استفاده از این ایده‌های تحلیلی در علوم مغزی نیز هست: بیش از ۱۰۰ مورد از چنین مثال‌هایی در سراسر کتاب وجود دارند و همگی در فهرست مثال‌ها آمده‌اند. خواننده‌ای که علاقه‌مند به این باشد که ببیند ما مثلاً چگونه به داده‌های fMRI پرداخته‌ایم، باید با فهرست مثال‌ها شروع کند. راهنماهای سازماندهی‌تر شده‌ای در فصل ۱ ارائه شده‌اند.

این کتاب یا به عنوان یک مرجع یا یک کتاب درسی طراحی شده است. آر. ای. کی از نسخه‌های اولیه دست‌نویس در کلاس‌هایی با حضور دانشجویان کارشناسی با پیشینه‌های متعدد از زیست‌شناسانی با دانش ریاضیاتی کم (که به دنبال درک مفاهیم بودند) تا

مهندسان (که به ریشه‌های مطالب نیاز داشتند) استفاده کرده بود. ما تصمیم گرفتیم که سعی کنیم هر دوی این مستمعین را راضی کنیم.

پیوستی شامل یک یادآور از ایده‌های ریاضیاتی کلیدی و ریشه‌های مطالب اغلب اختیاری در نظر گرفته می‌شوند. برای آن‌هایی که می‌خواهند از این کتاب به عنوان یک کتاب درسی استفاده کنند، آر. ای. کی مطالب را به ترتیب زیر پیشنهاد می‌کند:

• بخش I (آمار مقدماتی): فصل‌های ۱ تا ۷، ۱۰، ۱۲.۱ تا ۱۲.۴، ۱۳.۱.

• بخش II (نظریه آمار پایه): فصل‌های ۸، ۹، ۱۱، ۱۲.۵، ۱۳.۲ تا ۱۳.۴.

• بخش III (مباحث پیشرفته): گزیده‌هایی از فصل‌های ۱۴ تا ۱۹.

• به تجربه وی، بخش‌های I و II تقریباً، به ترتیب، ۱۲ و ۷ جلسه زمان می‌برند.

بسیاری از خوانندگان می‌خواهند کدهای رایانه‌ای برای روش‌هایی را که توضیح داده‌ایم ببینند. ما خودمان برای تولید شکل‌ها هم از MATLAB و هم از R استفاده کرده‌ایم. گرچه تصمیم گرفتیم که کدهای MATLAB یا R را داخل متن کتاب بیاوریم، کدها را در پایگاه اینترنتی مان به نشانی قرار داده‌ایم.

علاوه بر بسیاری از همکاران و دانشجویایی که پیشنهاداتی در طول مسیر ارائه کردند، شامل آن‌هایی که در متن از آن‌ها نام برده شده است، ما به اسپنسر کوثرنر مرهونیم که به ما در مرتب‌سازی و نوشتن بسیاری از کدها و تولید بسیاری از شکل‌ها کمک کرد، به پاتریک فولی که پایگاه اینترنتی کتاب را ساخت، هایدی سستریچ که بسیاری از نقص‌های فنی در LATEX را برطرف کرد، و متیو مارلر که همه متن پیش‌نویس را به‌دقت خواند و پیشنهادات به‌شدت سودمندی ارائه کرد. ما همچنین از آلن کوهن و رایان سایبرگ که هر یک چندین شکل از کتاب را تولید کردند سپاسگزاریم.

رابرت ای. کاس

اوری تی. ادن

امری ان. براون

فهرست کوتاه مطالب

ب	فهرست مطالب
۱	۱ مقدمات
۳۱	۲ کاوش در داده‌ها
۴۷	۳ احتمال و متغیرهای تصادفی
۹۱	۴ بردارهای تصادفی
۱۳۳	۵ توزیع‌های احتمال مهم
۱۷۳	۶ دنباله‌هایی از متغیرهای تصادفی
۱۸۷	۷ برآورد و عدم قطعیت
۲۲۵	۸ برآورد نظری و عملی
۲۷۵	۹ بوت‌استرپ و انتشار عدم قطعیت
۳۰۹	۱۰ مدل‌ها، فرض‌ها، و معنی‌داری آماری
۳۶۱	۱۱ روش‌های عمومی برای آزمودن فرض‌ها
۳۸۹	۱۲ رگرسیون خطی
۴۵۳	۱۳ تحلیل واریانس
۴۹۱	۱۴ رگرسیون خطی تعمیم‌یافته و رگرسیون غیرخطی
۵۱۹	۱۵ رگرسیون ناپارامتری
۵۵۱	۱۶ روش‌های بیزی
۶۱۵	۱۷ تحلیل چندمتغیری
۶۴۵	۱۸ سری‌های زمانی
۷۰۵	۱۹ فرایندهای نقطه‌ای
۷۶۱	پیوست‌ها
۷۸۳	مراجع
۷۹۳	واژه‌نامه

فهرست مطالب

ب	فهرست مطالب
۱	مقدمات
۱.۱	تحلیل داده‌ها در علوم مغز
۱.۱.۱	راهبردهای تحلیلی مناسب به‌شدت به هدف مطالعه و روش گردآوری داده‌ها بستگی دارند.
۲.۱.۱	بسیاری از بررسی‌ها شامل پاسخی به یک محرک یا رفتار هستند.
۲.۱	سهم آمار
۱.۲.۱	مدل‌های آماری نظم و تغییرپذیری داده‌ها را بر حسب توزیع‌های آماری توصیف می‌کنند.
۲.۲.۱	از مدل‌های آماری برای تبیین دانش و عدم قطعیت دربارهٔ سیگنال در حضور نویز از طریق استدلال استقرائی استفاده می‌شود.
۳.۲.۱	مدل‌های آماری ممکن است پارامتری یا ناپارامتری باشند.
۴.۲.۱	ساخت مدل آماری یک فرایند تکرارشونده همراه با ارزیابیِ برازش است و با تحلیل اکتشافی پیش می‌رود.
۵.۲.۱	همهٔ مدل‌ها نادرستند، اما برخی از آن‌ها مفید هستند.
۶.۲.۱	از نظریهٔ آمار برای درک رفتار شیوه‌های آماری تحت مفروضات احتمالاتی متعدد استفاده می‌شود.
۷.۲.۱	ایده‌های تحلیل داده‌ایِ مهم گاه به روش‌های متفاوتِ بسیاری پیاده‌سازی می‌شوند.
۸.۲.۱	دستگاه‌های اندازه‌گیری اغلب داده‌ها را پیش‌پردازش می‌کنند.
۹.۲.۱	فنون تحلیلی داده‌ها به‌ندرت قادر به جبرانِ نواقص موجود در گردآوری داده‌ها هستند.
۱۰.۲.۱	روش‌های ساده ضروری هستند.
۱۱.۲.۱	راحت است که داده‌ها را به چندین نوع گسترده طبقه‌بندی کنیم.
۳۱	کاوش در داده‌ها
۱.۲	تبیین گرایش به مرکز و پراکندگی

۱.۱.۲	نمودارها و خلاصه‌های مختلف چشم‌اندازهای متفاوتی از داده‌ها ارائه می‌کنند	۳۲
۲.۱.۲	روش‌های کاوشگرانه ممکن است پیچیده باشند	۳۶
۲.۲	تبدیلات داده‌ها	۳۸
۱.۲.۲	مقادیر مثبت اغلب با لگاریتم تبدیل می‌شوند	۳۸
۲.۲.۲	گاهی اوقات تبدیل‌های غیر لگاریتمی اعمال می‌شوند	۴۵
۴۷	۳ احتمال و متغیرهای تصادفی	
۱.۳	محاسبه احتمال	۴۸
۱.۱.۳	احتمال‌ها روی مجموعه‌هایی از پیشامدهای غیرقطعی تعریف می‌شوند	۴۸
۲.۱.۳	احتمال شرطی $P(A B)$ برابر است با احتمال این که A اتفاق می‌افتد به شرط آن که B اتفاق افتاده باشد	۵۱
۳.۱.۳	احتمال‌ها وقتی پیشامدها مستقل هستند در هم ضرب می‌شوند	۵۳
۴.۱.۳	قضیه بیز برای پیشامدها احتمال شرطی $P(A B)$ را برحسب احتمال شرطی $P(B A)$ به دست می‌دهد	۵۴
۲.۳	متغیرهای تصادفی	۵۹
۱.۲.۳	متغیرهای تصادفی مقادیری را می‌گیرند که با پیشامدها مشخص شده باشند	۶۱
۲.۲.۳	توزیع‌های متغیرهای تصادفی با استفاده از تابع‌های توزیع تجمعی و تابع‌های چگالی احتمال تعریف می‌شوند که از روی آن‌ها میانگین‌ها و واریانس‌های نظری را می‌توان محاسبه کرد	۶۲
۳.۲.۳	متغیرهای تصادفی پیوسته مشابه با متغیرهای تصادفی گسسته هستند	۶۷
۴.۲.۳	تابع خطر این احتمال شرطی را ارائه می‌دهد که پیشامدی رخ دهد، به شرط آن که تاکنون رخ نداده است	۷۷
۵.۲.۳	توزیع تابعی از یک متغیر تصادفی با فرمول‌های تغییر متغیر پیدا می‌شود	۷۸
۳.۳	تابع توزیع تجمعی تجربی	۸۲
۱.۳.۳	نمودارهای $P-P$ و $Q-Q$ کنترل‌هایی گرافیکی برای انحرافات زیاد از یک شکل توزیعی فراهم می‌کنند	۸۳
۲.۳.۳	از نمودارهای $Q-Q$ و $P-P$ می‌توان برای قضاوت درباره اثر بخشی تبدیل‌ها استفاده کرد	۸۸
۹۱	۴ بردارهای تصادفی	
۱.۴	دو یا چند متغیر تصادفی	۹۱
۱.۱.۴	تغییرات چندین متغیر تصادفی با توزیع توأم آن‌ها تبیین می‌شود	۹۳
۲.۱.۴	متغیرهای تصادفی، وقتی تابع چگالی توأم آن‌ها برابر با ضرب تابع‌های چگالی حاشیه‌ای است مستقل هستند	۹۶
۲.۴	وابستگی دو متغیری	۹۸

۱.۲.۴	وابستگی خطی دو متغیر تصادفی را می‌توان با همبستگی آن‌ها
۹۸	ارزیابی کرد.
۲.۲.۴	یک توزیع نرمال دو متغیری با یک جفت میانگین، یک جفت انحراف
۱۰۴	استاندارد و یک ضریب همبستگی تعیین می‌شود.
۳.۲.۴	احتمالات شرطی متغیرهای تصادفی را باید از چگالی‌های شرطی
۱۰۶	به‌دست آورد.
۴.۲.۴	امید ریاضی شرطی $E(Y X=x)$ رگرسیون Y روی X نامیده می‌شود. ...
۱۱۴	۳.۴ وابستگی چندمتغیری
۱۱۴	۱.۳.۴ میانگین یک بردار تصادفی، یک بردار و واریانس آن یک ماتریس است.
۱۱۶	۲.۳.۴ وابستگی دو بردار تصادفی را می‌توان با اطلاع متقابل ارزیابی کرد.
۱۲۳	۳.۳.۴ قضیهٔ بیز برای بردارهای تصادفی، مشابه با قضیهٔ بیز برای پیشامدها است.
۱۲۴	۴.۳.۴ رده‌بندهای بیزی بهینه‌اند.
۱۳۳	۵ توزیع‌های احتمال مهم
۱۳۴	۱.۵ متغیرهای تصادفی برنولی و توزیع دوجمله‌ای
۱۳۴	۱.۱.۵ متغیرهای تصادفی برنولی مقدارهای ۰ و ۱ را می‌گیرند.
۱۳۴	۲.۱.۵ توزیع دوجمله‌ای از جمع متغیرهای تصادفی برنولی مستقل و همگن
۱۳۴	حاصل می‌شود.
۱۴۰	۲.۵ توزیع پواسون
۱۴۰	۱.۲.۵ از توزیع پواسون اغلب در توصیف تعداد پیشامدهای دوحالتی استفاده
۱۴۰	می‌شود.
۱۴۴	۲.۲.۵ برای n بزرگ و p کوچک توزیع دوجمله‌ای تقریباً با پواسون یکسان
۱۴۴	است.
۱۴۶	۳.۲.۵ توزیع پواسون وقتی پیشامدهای دوحالتی مستقل هستند، حاصل
۱۴۶	می‌شود.
۱۴۹	۳.۵ توزیع نرمال
۱۴۹	۱.۳.۵ متغیرهای تصادفی نرمال با احتمال $\frac{1}{2}$ بین ۱ انحراف استاندارد از
۱۴۹	میانگین‌شان قرار دارند؛ با احتمال 0.95 بین ۲ انحراف استاندارد از
۱۴۹	میانگین‌شان قرار دارند.
۱۵۱	۲.۳.۵ توزیع‌های دوجمله‌ای و پواسون برای n بزرگ یا n بزرگ تقریباً
۱۵۱	نرمال هستند.
۱۵۱	۴.۵ برخی توزیع‌های متداول دیگر
۱۵۱	۱.۴.۵ توزیع چندجمله‌ای دوجمله‌ای را به چند رسته بسط می‌دهد.
۱۵۳	۲.۴.۵ از توزیع نمایی برای توصیف زمان‌های انتظار بی‌حافظه استفاده می‌شود.
۱۵۷	۳.۴.۵ توزیع‌های گاما و مجموع‌های نمایی‌ها
۱۵۸	۴.۴.۵ توزیع‌های خی‌دو حالت‌های خاصی از توزیع‌های گاما هستند.

۵.۴.۵	از توزیع بتا می‌توان برای توصیف پراکندگی بازه‌های متناهی استفاده کرد.	۱۵۸
۶.۴.۵	توزیع گوسی معکوس زمان انتظار برای یک آستانه گذرا با حرکت براونی را توصیف می‌کند.	۱۶۰
۷.۴.۵	توزیع‌های t و F از توزیع‌های نرمال و χ^2 تعریف می‌شوند.	۱۶۲
۵.۵	توزیع‌های نرمال چندمتغیری	۱۶۴
۱.۵.۵	یک بردار تصادفی نرمال چندمتغیری است اگر ترکیب‌های خطی مؤلفه‌های آن نرمال یک‌متغیری باشند.	۱۶۴
۲.۵.۵	تابع چگالی احتمال نرمال چندمتغیری دارای کانتورهای بیضوی با چگالی احتمال رو به افول طبق یک تابع چگالی احتمال است.	۱۶۵
۳.۵.۵	اگر X و Y دارای توزیع توأم نرمال چندمتغیری باشند، آن‌گاه توزیع شرطی Y به شرط X نرمال چندمتغیری است.	۱۶۸
۶	دنباله‌هایی از متغیرهای تصادفی	۱۷۳
۱.۶	دنباله‌های تصادفی و میانگین نمونه‌ای	۱۷۴
۱.۱.۶	انحراف استاندارد میانگین نمونه‌ای با ضریب $1/\sqrt{n}$ کاهش می‌یابد.	۱۷۶
۲.۱.۶	دنباله‌های تصادفی ممکن است طبق ملاک‌های مختلفی همگرا شوند.	۱۷۹
۲.۶	قانون اعداد بزرگ	۱۸۱
۱.۲.۶	با افزایش اندازه نمونه n ، میانگین نمونه‌ای به میانگین نظری همگرا می‌شود.	۱۸۱
۲.۲.۶	تابع توزیع تجمعی تجربی به تابع توزیع تجمعی نظری همگرا است.	۱۸۲
۳.۶	قضیه حد مرکزی	۱۸۳
۱.۳.۶	برای n های بزرگ، میانگین نمونه‌ای به‌طور تقریبی دارای توزیع نرمال است.	۱۸۳
۲.۳.۶	برای n های بزرگ، میانگین چندمتغیری به‌طور تقریبی توزیع نرمال چندمتغیری دارد.	۱۸۵
۷	برآورد و عدم قطعیت	۱۸۷
۱.۷	برازش مدل‌های آماری	۱۸۷
۲.۷	مسئله برآورد	۱۹۱
۱.۲.۷	روش گشتاوری از میانگین و واریانس نمونه‌ای برای برآورد میانگین و واریانس نظری استفاده می‌کند.	۱۹۲
۲.۲.۷	روش ماکسیمم درست‌نمایی تابع درست‌نمایی را، که تا یک ثابت ضریبی تعریف می‌شود، ماکسیمم می‌کند.	۱۹۳
۳.۷	بازه‌های اطمینان	۱۹۸
۱.۳.۷	برای استنباط‌های علمی، برآوردها بدون ذکر دقتشان بی‌فایده هستند.	۱۹۸
۲.۳.۷	برآورد یک میانگین نرمال، یک حالت پارادایم است.	۲۰۱

۳.۳.۷	برای مشاهدات غیر نرمال، می‌توان به قضیه حد مرکزی استناد کرد.	۲۰۲
۴.۳.۷	یک بازه اطمینان بزرگ‌نمونه‌ای برای μ با استفاده از خطای استاندارد	
۲۰۳	$s/\sqrt{n}$ به دست می‌آید.	
۵.۳.۷	خطاهای استاندارد بلافاصله به بازه‌های اطمینان می‌انجامند.	۲۰۵
۶.۳.۷	برآوردها و خطاهای استاندارد را باید با دو رقم اعشار از مقدار خطای	
۲۱۱	استاندارد گزارش کرد.	
۷.۳.۷	اندازه‌های نمونه تقریبی را می‌توان از روی اندازه مطلوب خطای	
۲۱۱	استاندارد تعیین کرد.	
۸.۳.۷	اطمینان، احتمال را به صورت غیر مستقیم تخصیص می‌دهد تا تفسیر	
۲۱۳	آن را دقیق کند.	
۹.۳.۷	از قضیه بیز می‌توان برای ارزیابی عدم قطعیت استفاده کرد.	۲۱۶
۱۰.۳.۷	برای نمونه‌های کوچک مرسوم است که از توزیع t به جای نرمال	
۲۲۱	استفاده شود.	
۲۲۵	۸ برآورد نظری و عملی	
۲۲۸	۱.۸ میانگین توان‌های دوم خطا	
۲۲۹	میانگین توان‌های دوم خطا همان مجذور اریبی به علاوه واریانس است.	
۲۱.۸	میانگین توان‌های دوم خطا را می‌توان با شبیه‌سازی رایانه‌ای	
۲۳۴	شبه‌داده‌ها محاسبه کرد.	
۳.۱.۸	در برآورد یک میانگین نظری از مشاهداتی که واریانس‌های متفاوتی	
	دارند، یک میانگین وزنی باید استفاده شود، که وزن‌ها به طور معکوس	
۲۳۹	متناسب با واریانس‌ها هستند.	
۴.۱.۸	نظریه تصمیم اغلب از میانگین توان دوم خطا برای بیان ریسک	
۲۴۶	استفاده می‌کند.	
۲۴۶	۲.۸ برآورد در نمونه‌های بزرگ	
۱.۲.۸	در نمونه‌های بزرگ، یک برآوردگر باید با احتمال زیاد نزدیک به آنچه	
۲۴۶	برآورد شده باشد.	
۲.۲.۸	در نمونه‌های بزرگ، کران دقت برآورد یک پارامتر برابر با اطلاع	
۲۴۷	فیشر است.	
۳.۲.۸	برآوردگرهایی که واریانس بزرگ نمونه‌ای را مینیمم می‌کنند کارا	
۲۵۲	نامیده می‌شوند.	
۲۵۴	۳.۸ ویژگی‌های برآوردگرهای ML	
۲۵۴	در نمونه‌های بزرگ، برآورد ML بهینه است.	
۲.۳.۸	خطای استاندارد MLE از مشتق دوم لگاریتم تابع درست‌نمایی	
۲۵۴	به دست می‌آید.	
۳.۳.۸	در نمونه‌های بزرگ، برآورد ML به طور تقریبی بیزی است.	۲۵۸
۴.۳.۸	برآوردهای ML همراه با پارامترها تبدیل می‌شوند.	۲۵۹

۵.۳.۸	ماکسیمم درست‌نمایی، با شرط نرمال بودن، میانگین وزنی را ارائه می‌دهد.	۲۶۰
۴.۸	ماکسیمم درست‌نمایی چندپارامتری	۲۶۱
۱.۴.۸	برآوردگر MLE از حل مجموعه‌ای از معادلات دیفرانسیل	
۲۶۱	جزییه‌دست می‌آید.	۲۶۱
۲.۴.۸	می‌توان کمترین توان‌های دوم را به عنوان حالتی خاص از برآورد	
۲۶۳	ML در نظر گرفت.	۲۶۳
۳.۴.۸	اطلاع مشاهده‌شده، منفی ماتریس مشتقات جزئی مرتبه دوم تابع	
۲۶۵	لگاریتم درست‌نمایی است که در $\hat{\theta}$ محاسبه شده است.	۲۶۵
۴.۴.۸	وقتی از روش‌های عددی برای برآورد ML استفاده می‌شود، کمی	
۲۶۷	احتیاط لازم است.	۲۶۷
۵.۴.۸	گاهی اوقات برآوردهای ML از الگوریتم EM به‌دست می‌آیند.	۲۶۸
۶.۴.۸	ماکسیمم درست‌نمایی ممکن است برآوردهایی بد تولید کند.	۲۷۳
۹	بوت‌استرپ و انتشار عدم قطعیت	۲۷۵
۱.۹	انتشار عدم قطعیت	۲۷۸
۱.۱.۹	مشاهدات شبیه‌سازی‌شده از توزیع متغیر تصادفی X مشاهدات	
۲۷۸	شبیه‌سازی‌شده از توزیع متغیر تصادفی $Y = f(X)$ را تولید می‌کنند.	۲۷۸
۲.۱.۹	در نمونه‌های بزرگ، تبدیل‌های متغیرهای تصادفی سازگار و مجانباً	
۲۸۵	نرمال تقریباً خطی خواهند بود.	۲۸۵
۲.۹	بوت‌استرپ	۲۹۵
۱.۲.۹	بوت‌استرپ روشی عمومی در ارزیابی عدم قطعیت است.	۲۹۶
۲.۲.۹	بوت‌استرپ پارامتری، شبه‌داده‌ها را از یک توزیع پارامتری برآورده‌شده	
۲۹۸	استخراج می‌کند.	۲۹۸
۳.۲.۹	بوت‌استرپ ناپارامتری شبه‌داده‌ها را از تابع توزیع تجمعی تجربی	
۳۰۱	استخراج می‌کنند.	۳۰۱
۳.۹	بحثی درباره روش‌های جایگزین	۳۰۶
۱۰	مدل‌ها، فرض‌ها، و معنی‌داری آماری	۳۰۹
۱.۱.۰	آماره‌های خی‌دو	۳۱۱
۱.۱.۱.۰	آماره خی‌دو مقادیر برازش یافته از مدل را با مقادیر مشاهده‌شده	
۳۱۲	مقایسه می‌کند.	۳۱۲
۲.۱.۱.۰	برای داده‌های چندجمله‌ای، آماره خی‌دو تقریباً دارای توزیع χ^2 است.	۳۱۴
۳.۱.۱.۰	نادربودن یک آزمون خی‌دوی بزرگ از طریق p -مقدارش مشخص	
۳۱۷	می‌شود.	۳۱۷
۴.۱.۱.۰	خی‌دو را می‌توان برای آزمون استقلال دو صفت به‌کار برد.	۳۱۹
۲.۱.۰	فرض‌های صفر	۳۲۲
۱.۲.۱.۰	اغلب مدل‌های آماری به عنوان فرض‌های صفر در نظر گرفته می‌شوند.	۳۲۲

۲.۲.۱۰	گاهی اوقات فرض‌های صفر مقدار خاصی از یک پارامتر را در یک
۳۲۲	مدل آماری مشخص می‌کنند.
۳.۲.۱۰	فرض‌های صفر همچنین می‌توانند قیدی روی دو یا چند پارامتر قرار
۳۲۳	دهند.
۳.۱۰	آزمون فرض‌های صفر
۳۲۴	
۱.۳.۱۰	آزمون $H_0 = \mu = \mu_0$ برای یک متغیر تصادفی نرمال، یک حالت
۳۲۴	پارادایم است.
۲.۳.۱۰	برای اندازه‌های نمونه بزرگ، فرض $H_0: \theta = \theta_0$ را می‌توان با استفاده
۳۲۶	از نسبت $(\hat{\theta} - \theta_0)/SE(\hat{\theta})$ آزمون کرد.
۳.۳.۱۰	برای نمونه‌های کوچک استفاده از آماره t برای آزمون مرسوم است. . .
۳۲۸	برای دو نمونه مستقل، فرض $H_0: \mu_1 = \mu_2$ را می‌توان با استفاده از
۳۳۰	نسبت t آزمون کرد.
۳۳۳	برای پیدا کردن p -مقادیر می‌توان از شبیه‌سازی رایانه‌ای استفاده کرد. .
۶.۳.۱۰	آزمون رایی می‌تواند شواهدی علیه دارای توزیع یکنواخت بودن
۳۳۶	زاویه‌ها ارائه کند.
۷.۳.۱۰	برآزش یک توزیع پیوسته را می‌توان با آزمون کولموگروف-اسمیرنوف
۳۳۸	ارزیابی کرد.
۴.۱۰	تفسیر و ویژگی‌های آزمون‌ها
۳۳۹	
۱.۴.۱۰	آزمون‌های آماری باید، حداقل به‌طور تقریبی، احتمال رد نادرست H_0
۳۴۰	صحیحی داشته باشند.
۲.۴.۱۰	یک بازه اطمینان برای θ را می‌توان برای آزمودن $H_0: \theta = \theta_0$ به‌کار برد. .
۳۴۴	آزمون‌های آماری بر حسب احتمال به‌درستی رد کردن H_0 ارزیابی
۳۴۶	می‌شوند.
۴.۴.۱۰	عملکرد یک آزمون آماری را می‌توان با منحنی ROC نمایش داد. . .
۳۴۸	
۵.۴.۱۰	p -مقدار احتمال این که H_0 درست باشد نیست
۳۵۲	
۶.۴.۱۰	p -مقدارهای مرزی واقعاً آزاردهنده و با توانی کم هستند.
۳۵۴	
۷.۴.۱۰	p -مقدار به‌طور مفهومی متمایز از خطای نوع اول است.
۳۵۵	
۸.۴.۱۰	آزمون غیر معنی دار، به خودی خود شواهدی در حمایت از H_0 نشان
۳۵۶	نمی‌دهد.
۹.۴.۱۰	گاهی اوقات آزمون‌های یک طرفه به‌کار گرفته می‌شوند.
۳۵۸	
۳۶۱	۱۱ روش‌های عمومی برای آزمودن فرض‌ها
۳۶۲	
۱.۱۱	آزمون‌های نسبت درست‌نمایی
۳۶۳	
۱.۱.۱۱	از نسبت درست‌نمایی می‌توان برای آزمودن $H_0: \theta = \theta_0$ استفاده کرد. .
۲.۱.۱۱	p -مقدارها برای آزمون نسبت درست‌نمایی $H_0: \theta = \theta_0$ را می‌توان از
۳۶۵	توزیع χ^2 یا با شبیه‌سازی به‌دست آورد

۳۶۷	۳.۱.۱۱ آزمون نسبت درست‌نمایی $(\omega, \theta) = (\omega, \theta_0)$: H_0 برآورد ML پارامتر ω را جای‌گذاری می‌کند و تحت H_0 به‌دست می‌آید.
۳۶۹	۴.۱.۱۱ آزمون نسبت درست‌نمایی، دقیقاً یا تقریباً، تعداد زیادی آزمون معنی‌داری متداول تولید می‌کند.
۳۶۹	۵.۱.۱۱ آزمون نسبت درست‌نمایی برای فرض‌های ساده بهینه است.
۳۷۱	۶.۱.۱۱ برای ارزیابی مدل‌های غیر آشیانه‌ای جایگزین می‌توان ابعاد پارامترهای آماره نسبت درست‌نمایی را تعدیل کرد.
۳۷۴	۲.۱۱ آزمون‌های جایگشتی و بوت‌استرپ
۳۷۴	۱.۲.۱۱ آزمون‌های جایگشتی همه جایگشت‌های ممکن داده‌هایی را که ممکن است با فرض صفر سازگار باشند در نظر می‌گیرند.
۳۷۸	۲.۲.۱۱ نمونه‌های بوت‌استرپ با جای‌گذاری.
۳۷۹	۳.۱۱ آزمون‌های چندگانه
۳۷۹	۱.۳.۱۱ وقتی از چند مجموعه داده مستقل برای آزمودن فرض‌های یکسان استفاده می‌شود، p -مقادیرها به راحتی ترکیب می‌شوند.
۳۸۱	۲.۳.۱۱ وقتی چند فرض در نظر گرفته می‌شوند، معنی‌داری آماری باید تعدیل شود.
۳۸۹	۱۲ رگرسیون خطی
۳۹۰	۱.۱۲ مدل رگرسیون خطی
۳۹۵	۱.۱.۱۲ رگرسیون خطی، خطی بودن $f(x)$ و استقلال اجزای نویز در مقادیر متفاوت مشاهده‌شده از x مفروض می‌داند.
۳۹۷	۲.۱.۱۲ سهم نسبی سیگنال خطی به کل تغییرپذیری پاسخ در R^2 خلاصه می‌شود.
۳۹۹	۳.۱.۱۲ نظریه نشان می‌دهد که اگر مدل درست باشد برآوردهای کمترین توان‌های دوم احتمالاً برای نمونه‌های بزرگ دقیق هستند.
۴۰۰	۲.۱۲ بررسی مفروضات
۴۰۰	۱.۲.۱۲ باقیمانده‌ها باید نویز بدون ساختار را نشان دهند.
۴۰۲	۲.۲.۱۲ بررسی گرافیکی داده‌های ممکن است اطلاعات مهمی را به‌دست دهد.
۴۰۳	۳.۲.۱۲ عدم استقلال خطاها ممکن است پیامدهای قابل توجهی داشته باشد.
۴۰۵	۳.۱۲ شواهد یک روند خطی
۴۰۵	۱.۳.۱۲ طبق فرمول کلی، بازه‌های اطمینان برای شیب‌ها براساس SE هستند.
۴۰۸	۲.۳.۱۲ شواهدی را که به نفع یک روند خطی هستند می‌توان با یک آزمون t در مورد شیب به‌دست آورد.
۴۰۹	۳.۳.۱۲ ممکن است رابطه خارج از محدوده داده‌های مشاهده شده دقیق نباشد.
۴۱۰	۴.۱۲ همبستگی و رگرسیون
۴۱۱	۱.۴.۱۲ ضریب همبستگی با ضریب رگرسیونی و انحراف استانداردهای x و y تعیین می‌شود.

۲.۴.۱۲	همبستگی، علّیت نیست	۴۱۱
۳.۴.۱۲	بازه‌های اطمینان برای p ممکن است بر مبنای یک تبدیل از r باشند.	۴۱۲
۴.۴.۱۲	هنگامی که نویز به دو متغیر اضافه می‌شود، همبستگی آن‌ها کاهش می‌یابد.	۴۱۴
۵.۱۲	رگرسیون خطی چندگانه	۴۱۶
۱.۵.۱۲	رگرسیون چندگانه، رابطه خطی پاسخ را با هر متغیر توضیحی برآورد می‌کند، در حالی که سایر متغیرهای توضیحی را تعدیل می‌کند.	۴۱۹
۲.۵.۱۲	تغییرات پاسخ را می‌توان به مجموع توان‌های دوم سیگنال و نویز تجزیه کرد.	۴۲۱
۳.۵.۱۲	رگرسیون چندگانه را می‌توان به‌طور مختصر با استفاده از ماتریس‌ها فرموله کرد.	۴۲۵
۴.۵.۱۲	مدل رگرسیون خطی در رگرسیون چندجمله‌ای و رگرسیون کسینوسی به کار می‌رود.	۴۳۳
۵.۵.۱۲	اثرات متغیرهای تبیینی همبسته را نمی‌توان به‌طور جداگانه تفسیر کرد.	۴۳۷
۶.۵.۱۲	برهم‌کنش‌های رگرسیون چندگانه اغلب مهم هستند.	۴۴۱
۷.۵.۱۲	اغلب می‌توان مدل‌های رگرسیونی با بسیاری از متغیرهای تبیینی را ساده کرد.	۴۴۱
۸.۵.۱۲	رگرسیون چندگانه ممکن است خیانتکار باشد.	۴۴۹
۴۵۳	۱۳ تحلیل واریانس	
۱.۱۳	آنوای یک‌راهه و دوراهه	۴۵۴
۱.۱.۱۳	آنوا مبتنی بر یک مدل خطی است.	۴۵۶
۲.۱.۱۳	آنوای دوراهه پراکندگی کل را به پراکندگی توسط گروه و پراکندگی متوسط فردی تجزیه می‌کند، که تحت فرض صفر تقریباً با هم برابرند.	۴۵۸
۳.۱.۱۳	وقتی فقط دو گروه وجود دارد، آزمون F در آنوا به یک آزمون t کاهش می‌یابد.	۴۶۲
۴.۱.۱۳	با آنوای دوراهه اثرات یک عامل را وقتی عامل دیگری تعدیل شده باشد ارزیابی می‌کنیم.	۴۶۳
۵.۱.۱۳	وقتی واریانس‌ها در بین شرایط همگن هستند می‌توان از یک آزمون نسبت درستنمایی استفاده کرد.	۴۶۵
۶.۱.۱۳	طرح‌های آزمایشی پیچیده‌تر را می‌توان با آنوا مطابقت داد.	۴۶۶
۷.۱.۱۳	تحلیل‌های اضافی، شامل مقایسه‌های چندگانه، ممکن است مستلزم تعدیل‌های p -مقدارها باشند.	۴۶۶
۲.۱۳	آنوا در قالب رگرسیون	۴۷۰
۱.۲.۱۳	مدل خطی عمومی در بردارنده مدل‌های رگرسیونی و مدل‌های آنوا است.	۴۷۰
۲.۲.۱۳	در آنوای چندراهه، برهم‌کنش‌ها اغلب مورد علاقه هستند.	۴۷۴

۳.۲.۱۳	مقایسه‌های آنوا را می‌توان با استفاده از تحلیل کوواریانس تعدیل کرد.	۴۷۷
۳.۱۳	روش‌های ناپارامتری	۴۷۹
۱.۳.۱۳	آزمون‌های ناپارامتری توزیع آزاد را می‌توان با جای‌گذاری مقادیر داده‌ای با رتبه‌های آن‌ها به‌دست آورد.	۴۸۰
۲.۳.۱۳	از آزمون‌های جایگشتی و بوت‌استرپ می‌توان برای آزمون فرض‌های آنوا استفاده کرد.	۴۸۳
۴.۱۳	علیت، تصادفی‌سازی، و مطالعات مشاهده‌ای	۴۸۴
۱.۴.۱۳	تصادفی‌سازی اثرات عامل‌های مزاحم را از بین می‌برد.	۴۸۴
۲.۴.۱۳	مطالعات مشاهده‌ای ممکن است شواهد قابل توجهی فراهم کنند.	۴۸۷
۱۴	رگرسیون خطی تعمیم‌یافته و رگرسیون غیرخطی	۴۹۱
۱.۱۴	رگرسیون لوژستیک، رگرسیون پواسونی و مدل‌های خطی تعمیم‌یافته	۴۹۳
۱.۱.۱۴	رگرسیون لوژستیک را می‌توان برای تحلیل پاسخ‌های دوحالتی به‌کار برد.	۴۹۳
۲.۱.۱۴	در رگرسیون لوژستیک، روش ماکسیمم درست‌نمایی برای برآورد ضرایب رگرسیونی به‌کار می‌رود و آزمون نسبت درست‌نمایی برای ارزیابی شواهد روند خطی-لوژستیکی با x به‌کار گرفته می‌شود.	۴۹۷
۳.۱.۱۴	تبدیل لوجیت یکی از تبدیلات بسیاری است که می‌توان برای پاسخ‌های دو جمله‌ای استفاده کرد، اما بیشترین کاربرد را دارد.	۵۰۰
۴.۱.۱۴	مدل رگرسیونی پواسون معمولی میانگین λ را به $\log \lambda$ تبدیل می‌کند.	۵۰۲
۵.۱.۱۴	در رگرسیون پواسونی، برای برآورد ضرایب از روش ماکسیمم درست‌نمایی و برای بررسی روندها از آزمون نسبت درست‌نمایی استفاده می‌شود.	۵۰۴
۶.۱.۱۴	مدل‌های خطی تعمیم‌یافته روش‌های رگرسیونی را به پاسخ‌هایی با توزیعی از خانواده نمایی تعمیم می‌دهند.	۵۰۵
۲.۱۴	رگرسیون غیرخطی	۵۰۸
۱.۲.۱۴	مدل‌های رگرسیون غیرخطی را می‌توان با کمترین توان‌های دوم برازش داد.	۵۰۸
۲.۲.۱۴	مدل‌های غیرخطی تعمیم‌یافته را می‌توان با استفاده از ماکسیمم درست‌نمایی برازش داد.	۵۱۳
۳.۲.۱۴	در حل مسائل بهینه‌سازی غیرخطی، مقادیر نقاط آغازین با اهمیت هستند و پارامتربندی مجدد ممکن است مفید باشد.	۵۱۷
۱۵	رگرسیون ناپارامتری	۵۱۹
۱.۱۵	هموارسازها	۵۲۱
۱.۱.۱۵	هموارسازهای خطی سریع هستند.	۵۲۲

۲۰.۱۵	برای هموارسازهای خطی، مقادیر تابع‌های برازش‌یافته از طریق «ماتریس کلاه‌دار» به‌دست می‌آیند و به‌کار بستن انتشار عدم قطعیت ساده است.	۵۲۲
۲۰.۱۵	تابع‌های پایه	۵۲۳
۱.۲.۱۵	می‌توان از اسپلاین‌ها برای نمایش تابع‌های پیچیده استفاده کرد.	۵۲۵
۲.۲.۱۵	اسپلاین را می‌توان با استفاده از مدل‌های خطی به داده‌ها برازش داد.	۵۲۶
۳.۲.۱۵	استفاده از اسپلاین‌ها در مدل‌های خطی تعمیم‌یافته نیز آسان است.	۵۲۹
۴.۲.۱۵	اسپلاین‌های رگرسیون، همواربودن برازش را تعداد و مکان گره‌ها کنترل می‌کنند.	۵۳۱
۵.۲.۱۵	اسپلاین‌های هموارکننده اسپلاین‌هایی با گره‌هایی در هر یک از x_i اما با ضرایبی کاهش‌یافته هستند که از روش ML تاوانیده به‌دست آمده‌اند.	۵۳۱
۶.۲.۱۵	روشی موسوم به BARS مجموعه‌گره‌ها را به‌صورت خودکار طبق یک معیار بیزی انتخاب می‌کند.	۵۳۳
۷.۲.۱۵	از هموارسازی اسپلاینی می‌توان همراه با چندین متغیر کمکی استفاده کرد.	۵۳۴
۸.۲.۱۵	اغلب از جایگزین‌های اسپلاین‌ها در رگرسیون ناپارامتری استفاده می‌شود.	۵۳۶
۳.۱۵	برازش محلی	۵۳۹
۱.۳.۱۵	رگرسیون هسته‌ای را با یک میانگین موزون تعریف‌شده با یک تابع چگالی احتمال برآورد می‌کند.	۵۴۰
۲.۳.۱۵	رگرسیون چندجمله‌ای محلی یک مسئله کمترین توان‌های دوم با وزن‌های تعریف‌شده با یک هسته را حل می‌کند.	۵۴۳
۳.۳.۱۵	ملاحظات نظری به توصیه‌هایی دربارهٔ پهنای باند برای هموارسازهای خطی می‌انجامد.	۵۴۴
۴.۱۵	برآورد چگالی	۵۴۵
۱.۴.۱۵	از هسته‌ها می‌توان برای برآورد یک تابع چگالی احتمال استفاده کرد.	۵۴۶
۲.۴.۱۵	از روش‌های رگرسیون ناپارامتری دیگری می‌توان برای برآورد یک تابع چگالی احتمال استفاده کرد.	۵۴۸
۱۶	روش‌های بیزی	۵۵۱
۱.۱۶	توزیع‌های پسین	۵۵۳
۱.۱.۱۶	استنباط بیزی احتمال توصیفی و معرفت‌شناسانه را برابر در نظر می‌گیرد.	۵۵۴
۲.۱.۱۶	کار با پیشین‌های مزدوج راحت است.	۵۵۵
۳.۱.۱۶	برای خانواده‌های نمایی با پیشین‌های مزدوج، میانگین پسین، ترکیبی وزنی از MLE و میانگین پیشین است.	۵۵۷

۴.۱.۱۶	انتخاب توجیه‌پذیری برای توزیع پیشین وجود ندارد.	۵۶۱
۵.۱.۱۶	برای نمونه‌های بزرگ، پسین‌ها تقریباً نرمال و در MLE متمرکز هستند.	۵۶۳
۶.۱.۱۶	روش‌های قدرتمندی برای محاسبه توزیع پسین وجود دارد.	۵۶۵
۲.۱۶	متغیرهای پنهان	۵۷۴
۱.۲.۱۶	مدل‌های سلسله مراتبی برآوردهایی از کمیت‌های مرتبطی را که به سمت یکدیگر کشیده می‌شوند، تولید می‌کنند.	۵۷۶
۲.۲.۱۶	برای مدل‌های سلسله‌مراتبی، توزیع پسین اغلب با استفاده از نمونه‌گیری گیبس محاسبه می‌شود.	۵۸۷
۳.۲.۱۶	رگرسیون توانیده را می‌توان به عنوان برآورد بیزی در نظر گرفت.	۵۹۰
۴.۲.۱۶	مدل‌های حالت فضا به پارامترها اجازه می‌دهند که به صورت پویا تکامل پیدا کنند.	۵۹۱
۵.۲.۱۶	از فیلتر کالمن می‌توان برای برآورد متغیرهای حالت در مدل‌های خطی حالت-فضای گوسی استفاده کرد.	۵۹۵
۳.۱۶	عامل‌های بیزی	۵۹۹
۱.۳.۱۶	عامل‌های بیز می‌توانند شواهد را به نفع فرض‌ها ارائه دهند.	۶۰۰
۲.۳.۱۶	عامل‌های بیزی تفسیری از پیشرفت علمی ارائه می‌دهند.	۶۰۴
۳.۳.۱۶	وقتی که اطلاعات اندکی در مورد پارامترهای نامعلوم وجود دارد، استفاده از عامل‌های بیز ممکن است دشوار باشد.	۶۰۵
۴.۳.۱۶	عامل‌های بیز را می‌توان برای کالبدن p -مقدارها استفاده کرد.	۶۰۷
۴.۱۶	نتیجه‌گیری درباره متغیرهای پنهان	۶۰۸

۱۷ تحلیل چندمتغیری

۱.۱۷	مقدمه	۶۱۵
۲.۱۷	تحلیل واریانس چندمتغیری	۶۱۷
۱.۲.۱۷	مانوا یک بسط چندمتغیری از آنوا است.	۶۱۷
۲.۲.۱۷	وقتی ماتریس‌های واریانس بین شرایط، یکسان نیستند، می‌توان آزمون نسبت درستنمایی را به کار برد.	۶۲۱
۳.۱۷	کاهش بعد	۶۲۴
۱.۳.۱۷	یک ماتریس واریانس را می‌توان به مؤلفه‌های اصلی تجزیه کرد.	۶۲۴
۲.۳.۱۷	می‌توان از روش‌هایی به جز PCA برای کاهش بعد استفاده کرد.	۶۳۰
۴.۱۷	رده‌بندی و خوشه‌بندی	۶۳۴
۱.۴.۱۷	تفکیک‌کننده‌های بیزی برای توزیع‌های نرمال چندمتغیری شکل ساده‌ای می‌گیرند.	۶۳۴
۲.۴.۱۷	تفکیک‌کننده‌های بیزی همیشه عملی نیستند.	۶۳۵
۳.۴.۱۷	مشاهدات چندمتغیری را می‌توان به گروه‌هایی خوشه‌بندی کرد.	۶۴۰

۱۸ سری‌های زمانی

۶۴۵	۱.۱۸ مقدمه
۶۵۲	۲.۱۸ زمان دامنه‌ای و فرکانس دامنه‌ای
۶۵۶	۱.۲.۱۸ تحلیل فوریه یکی از دستاوردهای بزرگ علوم ریاضی است.
	۲.۲.۱۸ دوره‌نگار هم نمایش مقیاس‌بندی شده از سهم R^2 از رگرسیون
	هارمونیک است و هم یک تابع توان مقیاس‌بندی شده مرتبط با تبدیل
۶۶۱	گسسته فوریه از یک مجموعه داده است.
۶۶۶	۳.۲.۱۸ مدل‌های اتورگرسیو را می‌توان با رگرسیون تأخیریافته برازش داد.
۶۷۲	۳.۱۸ دوره‌نگار فرایندهای مانا.
۶۷۲	۱.۳.۱۸ دوره‌نگار را می‌توان به عنوان برآوردی از تابع چگالی طیفی در نظر گرفت.
	۲.۳.۱۸ برای نمونه‌های بزرگ، عرض‌های دوره‌نگار محاسبه شده از یک سری
۶۷۴	زمانی مانا تقریباً مستقل از یکدیگر و دارای توزیع χ^2 هستند.
	۳.۳.۱۸ برآوردگرهای سازگار تابع چگالی طیفی از هموارسازی دوره‌نگار
۶۷۶	به دست می‌آیند.
۶۷۹	۴.۳.۱۸ فیلترهای خطی ممکن است سریع و مؤثر باشند.
۶۸۳	۵.۳.۱۸ اطلاعات فرکانس محدود به نرخ نمونه‌گیری است.
	۶.۳.۱۸ مخروطی شدن، نشتی را از فرکانس‌های غیر فوریه به فوریه کاهش
۶۸۳	می‌دهد.
۶۸۷	۷.۳.۱۸ تحلیل زمان-فرکانس تحول ریتم‌ها را در طول زمان بیان می‌کند.
۶۹۰	۴.۱۸ انتشار عدم قطعیت برای توابع دوره‌نگار
	۱.۴.۱۸ بازه‌های اطمینان و آزمون‌های معنی‌داری را می‌توان با انتشار عدم
۶۹۰	قطعیت از دوره‌نگار به دست آورد.
	۲.۴.۱۸ عدم قطعیت در مورد توابعی از سری‌های زمانی را می‌توان از
۶۹۲	شبه داده‌های سری زمانی به دست آورد.
۶۹۳	۵.۱۸ سری‌های زمانی دومتغیری
	۱.۵.۱۸ انسجام $\rho_{XY}(\omega)$ بین دو سری X و Y را می‌توان همبستگی
۶۹۵	مؤلفه‌های فرکانس ω آن‌ها در نظر گرفت.
	۲.۵.۱۸ در بررسی همبستگی متقابل یا انسجام دو سری زمانی، در ابتدا
۶۹۹	سفید کردن سری توصیه می‌شود.
	۳.۵.۱۸ علیت گرنجر، پیش‌بینی‌پذیری خطی یک سری زمانی توسط سری
۷۰۰	دیگر را اندازه‌گیری می‌کند.
۷۰۵	۱۹ فرایندهای نقطه‌ای
۷۱۰	۱.۱۹ نمایش‌های فرایندهای نقطه‌ای
	۱.۱.۱۹ یک فرایند نقطه‌ای را می‌توان بر حسب زمان‌های پیشامد، بازه‌های
۷۱۰	بین پیشامدی، یا تعداد پیشامدها مشخص کرد.
	۲.۱.۱۹ یک فرایند نقطه‌ای را می‌توان به‌طور تقریبی به عنوان یک سری
۷۱۲	زمانی دوحالتی در نظر گرفت.

۳.۱.۱۹	فرایندهای نقطه‌ای می‌توانند گستره وسیعی از رفتارهای وابسته به پیشینه را نمایش دهند.	۷۱۳
۲.۱۹	فرایندهای پواسون	۷۱۵
۱.۲.۱۹	فرایندهای پواسون فرایندهای نقطه‌ای هستند که برای آن‌ها احتمال‌های پیشامد به وقوع یا زمان پیشامدهای گذشته بستگی ندارد.	۷۱۵
۲.۲.۱۹	فرایندهای پواسون ناهمگن دارای بازه‌هایی متغیر با زمانند.	۷۲۰
۳.۱۹	فرایندهای نقطه‌ای ناپواسونی	۷۲۷
۱.۳.۱۹	فرایندهای تجدید دارای زمان‌های انتظار درون‌پیشامدی i.i.d. هستند.	۷۲۷
۲.۳.۱۹	تابع شدت شرطی چگالی احتمال توأم زمان‌های انتظار برای یک فرایند نقطه‌ای عمومی را مشخص می‌کنند	۷۳۲
۳.۳.۱۹	شدت حاشیه‌ای، امید ریاضی شدت شرطی است.	۷۳۵
۴.۳.۱۹	تابع‌های شدت شرطی را می‌توان با استفاده از رگرسیون پواسونی برازش داد.	۷۳۷
۵.۳.۱۹	کنترل‌های گرافیکی برای انحراف‌ها از یک مدل فرایند نقطه‌ای را می‌توان با استفاده از مقیاس‌بندی زمانی به‌دست آورد.	۷۴۸
۶.۳.۱۹	روش‌هایی کارا برای تولید شبه‌داده‌های فرایند نقطه‌ای وجود دارند.	۷۵۲
۷.۳.۱۹	تحلیل طیفی فرایندهای نقطه‌ای نیازمند احتیاط است	۷۵۴
۴.۱۹	استخراج‌هایی دیگر	۷۵۶
۷۶۱	پیوست‌ها	
۷۸۳	مراجع	
۷۹۳	واژه‌نامه	