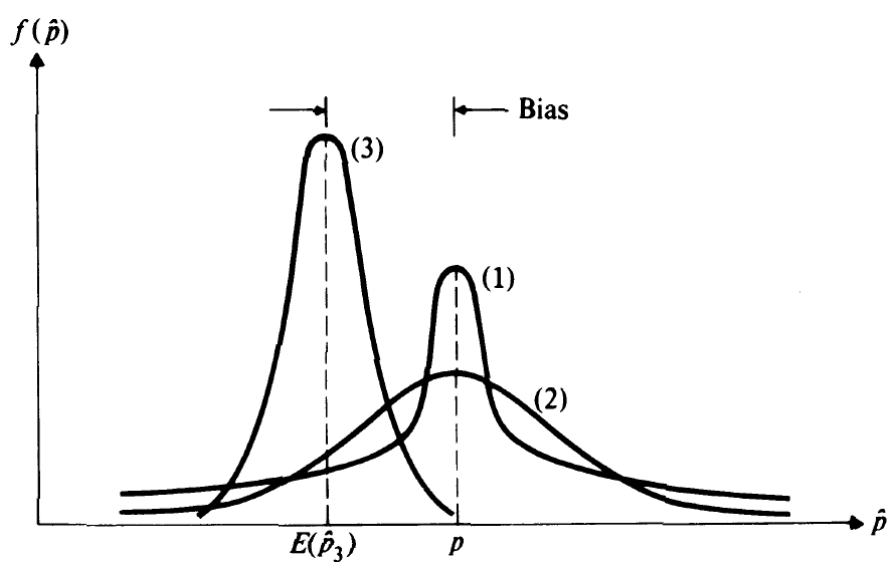


نظریه خطاها در مهندسی نقشه برداری



تهیه کننده : یحیی جمور

زمستان ۱۳۸۸

پیشگفتار

بشر از ابتدای خلقت تاکنون به دنبال کشف حقایق بوده و هست. بخش قابل توجهی از این حقایق جنبه فیزیکی دارند و دارای بعد زمان و مکان هستند. زمین به عنوان محل سکونت و زندگی انسان با شکل و اندازه خود یکی از موضوعاتی است که شناخت دقیقتر آن همواره مورد خواست بشر بوده است. شناخت ویژگی ها و خصوصیات زمین و محیط پیرامون زندگی بشر مبتنی بر دانستن یکسری کمیت های هندسی و فیزیکی است که مجهول می باشند. رشته مهندسی نقشه برداری در میان رشته های مختلف مرتبط با علوم زمین از جایگاه ممتازی در دستیابی به این شناخت برخوردار است. اندازه گیری های مختلف در حوزه زمان و مکان یک نقش زیربنایی در مسائل مهندسی نقشه برداری دارند. البته بدیهی است در سایر رشته های علوم و مهندسی نیز برای دستیابی به اهداف مورد نظرشان، قطعاً اندازه گیری ها از اهمیت ویژه ای برخوردارند.

هر چند امروزه با اتکا بر فناوری های جدید در مهندسی نقشه برداری نظیر سامانه های GNSS یک تحول اساسی در سرعت و دقت تعیین کمیت های مجهول به وجود آمده است، لیکن بسته به هدف و دقت مورد نظر همواره یک عدم قطعیت در برآورد کمیت های مجهول وجود دارد. ریشه این عدم قطعیت به وجود خطا در اندازه گیری ها و نیز مدل های مورد استفاده در مسائل مختلف مهندسی نقشه برداری بر می گردد. با توجه به اینکه هیچگاه امکان دستیابی به اندازه گیری های بدون خطا و مدلسازی های کامل و بی نقص وجود ندارد، لذا بررسی رفتار خطاها در اندازه گیری ها و سپس در مدلسازی ها و نتایج حاصل از آن، بویژه در مسائل حساس و دقیق، امری ضروری است که باید به طور جدی مد نظر قرار گیرد. بر این اساس یکی از مباحث مهم در مسائلی که به نوعی با اندازه گیری سروکار دارند "تئوری خطاها" است که همیشه به عنوان یکی از دروس پایه و مهم در رشته مهندسی نقشه برداری به حساب می آید.

بر اساس تجارب نگارنده جزوه پس از سال ها آموزش درس تئوری خطاها در دانشگاه ها و مراکز مختلف آموزش عالی و توجه به محتوای کتب موجود در زمینه تئوری خطاها و اندازه گیری در مهندسی نقشه برداری، مطالب زیر را در ۵ فصل برای استفاده دانشجویان عزیز تدوین شده است.

- مفهوم اندازه گیری و خطا
- مفاهیم اساسی آمار و احتمال
- انتشار خطاها
- نمونه برداری و برآورد
- بیضی خطاها

نگارنده امیدوار است این جزوه موجب رضایت کلیه خوانندگان محترم بویژه دانشجویان عزیز قرار بگیرد و با این وسیله توانسته باشد گامی هر چند کوچک در جهت توسعه و اعتلای مهندسی نقشه برداری به علاقمندان و دانش پژوهان و جامعه دانشگاهی رشته مهندسی نقشه برداری برداشته باشد.

دکتر یحیی جمور

دیماه ۱۳۸۸

فهرست مطالب

۱	پیشگفتار
۷	۱- اندازه گیری و خطا
۷	۱-۱- فرآیندهای قطعی و اتفاقی
۸	۲-۱- مفهوم اندازه گیری و خطا
۱۱	۳-۱- انواع خطاها
۱۲	۱-۳-۱- خطاهای بزرگ یا فاحش
۱۳	۲-۳-۱- خطاهای سیستماتیک
۱۵	۳-۳-۱- خطاهای اتفاقی
۱۶	۴-۱- ارقام معنی دار و گرد کردن اعداد
۱۹	۲- مروری بر مفاهیم اساسی آمار و احتمال
۱۹	۱-۲- توزیع فراوانی
۲۳	۲-۲- قابلیت اطمینان اندازه گیری ها
۲۴	۳-۲- فضای نمونه و پیشامد
۲۵	۴-۲- احتمال
۲۶	۵-۲- قوانین و شرایط مدل احتمال
۲۷	۶-۲- متغیرهای تصادفی
۲۷	۷-۲- توابع توزیع احتمال گسسته
۳۰	۸-۲- توابع توزیع احتمال پیوسته
۳۱	۱-۸-۲- تابع توزیع یکنواخت
۳۳	۲-۸-۲- تابع توزیع نرمال

۳۶	۲-۸-۳- تابع توزیع کای اسکور
۳۷	۲-۸-۴- توزیع احتمال t-student
۳۸	۲-۸-۵- توزیع احتمال F
۴۰	۲-۹-۹- توزیع احتمال چند متغیره
۴۱	۲-۹-۱- تابع چگالی نرمال چند متغیره
۴۳	۲-۱۰-۱- امید ریاضی
۴۴	۲-۱۰-۱- کوریانس و ضریب همبستگی
۴۵	۲-۱۰-۲- امید ریاضی توابع
۴۵	۲-۱۰-۲-۱- مدل خطی
۴۶	۲-۱۰-۲-۲- مدل غیر خطی
۴۷	۲-۱۱-۱- ماتریس های وریانس-کوریانس
۵۰	۳- انتشار خطاها
۵۰	۳-۱- انتشار تابع توزیع
۵۲	۳-۲- انتشار میانگین
۵۳	۳-۳- انتشار خطا های سیستماتیک
۵۶	۳-۴- انتشار خطا های اتفاقی (انتشار وریانس کوریانس)
۶۰	۳-۴-۱- انتشار وریانس کوریانس در مدل های خطی
۶۲	۳-۴-۲- انتشار وریانس کوریانس در مدل های غیر خطی
۶۶	۳-۵- انتشار ماتریس کوریانس (مبحث تکمیلی)
۷۲	۴- نمونه برداری و برآورد
۷۲	۴-۱- جمعیت و نمونه

۷۲	۲-۴- برآورد
۷۳	۳-۴- معیار های کیفیت برآوردگرها
۷۳	۱-۳-۴- نااریبی
۷۴	۲-۳-۴- سازگاری
۷۵	۳-۳-۴- کمترین وریانس و کارایی
۷۵	۴-۴- روش های برآورد
۷۵	۱-۴-۴- روش گشتاور ها
۷۶	۲-۴-۴- روش حداکثر درست نمایی
۷۶	۳-۴-۴- روش کمترین مربعات
۷۷	۵-۴- تعیین فواصل اطمینان (برآورد فاصله ای)
۷۸	۶-۴- میانگین وزن دار
۸۵	۵- نمایش هندسی پخش خطا ها
۸۵	۱-۵- بیضی خطا
۹۵	۲-۵- منحنی پدال
۹۷	۶- پیوست ها: جداول آماری

فصل اول

اندازه گیری و خطا

علیرغم پیشرفت علوم و فنون مختلف و بکارگیری دستگاههای جدید اندازه گیری، هنوز اندازه گیری ها و مشاهدات مختلف نقشه برداری عاری از خطا نمی باشند. در واقع به طور کلی همه اندازه گیری ها تا اندازه ای غیر دقیق اند و مقدار واقعی کمیت های فیزیکی مانند طول، زاویه و اختلاف ارتفاع را نمی توان یافت. گذشته از کامل نبودن دستگاههای اندازه گیری عوامل محیطی و انسانی نیز در واقعی نبودن اندازه گیری ها دخیل هستند. از آنجاکه مقدار واقعی از نظر ما قابل حصول نیست تلاش خواهیم کرد تا چگونگی یافتن دقیقترین مقدار، که به کمک مجموعه ای از اندازه گیری ها مشخص می شود، و چگونگی برآورد دقت این مقدار را نشان دهیم.

البته همیشه هدف از اندازه گیری ها لزوما کمینه سازی خطاهای اندازه گیری و رسیدن به بالاترین دقت نیست. نتیجه کم دقت نیز ممکن است به خوبی پاسخگوی نیاز مسئله باشد و خطای باقیمانده اثری بر روی استنتاجات تحلیلی نداشته باشد. به هر حال هر اندازه گیری یا مشاهده با مقدار واقعی آن کمیت یک اختلاف دارد که آن را خطای مشاهده می نامند و در این کتاب به بحث و بررسی در مورد آن خواهیم پرداخت. خطاهای مشاهداتی از قانون ساده ای پیروی نمی کنند و به طور کلی از علل متعددی ناشی می شوند. حتی در تکرار یک اندازه گیری با یک وسیله واحد به نتایج دقیقا یکسانی نمی رسیم و تغییرات هر چند کوچکی در آنها می بینیم. صرفنظر از اشتباهات، خطاهای اندازه گیری معمولا به دو دسته سیستماتیک و اتفاقی تقسیم می شوند. البته برخی اوقات جداسازی این دو دسته خطا مشکل است و در واقع خطاهای اندازه گیری موجود ترکیبی از آنها هستند.

فرآیندهای قطعی و اتفاقی

در عمل بسیاری از پدیده های فیزیکی وجود دارند که داده های تولیدی آنها با دقت قابل قبولی با مدل های ریاضی صریح قابل بیان هستند. این نوع پدیده ها موسوم به پدیده های قطعی هستند که در دو دسته کلی متناوب و نا متناوب تقسیم می شوند و هریک از اینها نیز در زیر دسته های دیگری تقسیم بندی می

شوند (Bendat & Piersol, 1971). حرکت ماهواره در یک مدار مشخص به دور زمین و یا سقوط آزاد یک جسم در خلاء، مثال‌های مناسبی برای پدیده‌های فیزیکی قطعی هستند. در مقابل پدیده‌های دیگری نیز وجود دارند که داده‌هایی با ماهیت قطعی تولید نمی‌کنند و رفتار آنها به هیچ وجه قابل مدلسازی و قابل پیش‌بینی نیست. در چنین فرآیندهایی هر مشاهده کاملاً مستقل از مشاهدات دیگر می‌باشد و در واقع یکی از بی‌نهایت حالتی است که می‌تواند رخ دهد. تغییر ارتفاع امواج در یک دریای خروشان را می‌توان به عنوان یک مثال خوب برای این نوع پدیده‌ها در نظر گرفت. همانطور که اشاره شد، در فرآیندهای قطعی می‌توان رفتار پدیده‌های فیزیکی را بر اساس مدل‌های ریاضی برای آینده پیش‌بینی نمود در حالیکه در فرآیندهای اتفاقی چنین چیزی امکان‌پذیر نیست و باید از علم احتمال و آمار برای تجزیه و تحلیل آنها استفاده نمود.

آنچه که ما در این مبحث به دنبال آن هستیم فرآیندهای اتفاقی است که خود به دو دسته ایستا (stationary) و نا ایستا (nonstationary) تقسیم می‌شوند. منظور از فرآیندهای ایستا فرآیندهای با خواص آماری ثابت در طول زمان می‌باشد (مانند وریانس و میانگین). بر عکس در فرآیندهای نا ایستا خواص آماری آنها با گذشت زمان ثابت نخواهند ماند (Bendat & Piersol, 1971). خوشبختانه اغلب پدیده‌هایی که ما با آنها سروکار داریم یا واقعاً ایستا هستند و یا اینکه فرض ایستایی با تقریب بسیار خوبی برای آنها اعتبار دارد.

مفهوم اندازه‌گیری و خطا

اگرچه اغلب اوقات مهندسين نقشه بردار با اندازه‌گیرها و سپس سرشکنی و تجزیه و تحلیل آنها سروکار دارند، ولی دامنه مسئولیت هر مهندس نقشه بردار در ارتباط با هر پروژه نقشه برداری متغیر و گسترده‌تر از آن است و از تجزیه و تحلیل اولیه تا ارائه نتایج به نگارهای مختلف را در بر می‌گیرد. بنابراین چنانچه یک مهندس نقشه بردار بخواهد وظیفه خود را به نحو مطلوب انجام دهد و هوشمندانه نسبت به جمع‌آوری، سرشکنی و تجزیه و تحلیل اندازه‌گیری‌ها اقدام نماید، لازم است فرآیند اندازه‌گیری را خوب بشناسد.

در اغلب اوقات افراد مختلف حتی متخصصین اندازه‌گیری‌ها نیز واژه‌های "اندازه‌گیری" و "مشاهده" با مفهوم یکسانی بکار می‌برند. اما در اصل باید توجه نمود که انجام هر اندازه‌گیری نیازمند یکسری مشاهدات است و هیچ اندازه‌ای حاصل نمی‌شود مگر قبل از آن چیزی مشاهده گردد. بنابراین عمل

"مشاهده" مقدم بر "اندازه گیری" است و هر "اندازه گیری" حاصل حداقل یک "مشاهده" می باشد. از آنجا که در مهندسی نقشه برداری و برخی علوم و فنون دیگر وابستگی شدیدی بین مشاهدات و اندازه گیری ها وجود دارد، عبارات "اندازه گیری" و "مشاهده" اغلب به صورت معادل و مترادف استعمال می شوند.

ممکن است اندازه گیری در ذهن ما به صورت یک عمل منفرد تصور شود، ولی حقیقت موضوع کمی پیچیده تر از این است و ممکن است در یک نقشه برداری دقیق هر اندازه گیری شامل چندین مرحله اساسی نظیر استقرار (سانتراژ)، نشانه روی، تطابق و قرائت باشد. با این حال در پایان تمام این مراحل تنها یک مقدار عددی برای بیان "اندازه گیری" یا "مشاهده" کمیت مورد نظر ارائه می شود. برای مثال عمل نسبتاً ساده اندازه گیری یک فاصله ۲۰ متری بین دو نقطه را با استفاده از یک متر نواری فولادی ۳۰ متری در نظر بگیرید. برای اندازه گیری دقیق فاصله بین این دو نقطه مراحل اساسی زیر باید انجام شوند:

۱- نوار متر در نزدیکی دو سر ابتدایی و انتهایی فاصله مورد نظر، بالای سطح زمین و تقریباً در امتداد مستقیم بین دو نقطه نگهداشته می شود.

۲- روی نقطه اول و در مجاورت نوار متر، یک شاقول به حالت آویزان نگهداشته می شود.

۳- روی نقطه دوم و در مجاورت نوار متر، یک شاقول به حالت آویزان نگهداشته می شود.

۴- در حالیکه هر دو شاقول روی دو نقطه حفظ می شوند، متر به حالت کشیده در می آید.

۵- عدد مجاور نخ شاقول اول از روی متر قرائت و ثبت می شود.

۶- عدد مجاور نخ شاقول دوم از روی متر قرائت و ثبت می شود.

۷- با کم کردن دو قرائت مربوط به مراحل ۵ و ۶ فاصله بین دو نقطه بدست می آید.

ملاحظه می شود که برای چنین اندازه گیری ساده ای چندین مرحله اساسی شامل استقرار، کشش متر و قرائت انجام می گیرد که همه آنها برای بدست آوردن فاصله دقیق بین دو نقطه لازم و ضروری هستند. بنابراین اندازه گیری فاصله بین دو نقطه با روش متر کشی فقط یک قرائت واحد نیست و در واقع نتیجه چندین مرحله در فرآیند مشاهده و اندازه گیری کمیت مورد نظر است.

مراحل ۱ تا ۷ ممکن است هنوز تمام فرایند مربوط به بدست آوردن یک اندازه گیری قابل اعتماد از

فاصله مورد نظر نباشند. به عبارت دیگر ممکن است برای خیلی از مقاصد، مقدار عددی حاصل از مرحله ۷

مناسب باشد، اما برای مقاصدی دیگر اینطور نباشد. چنانچه نتیجه بهتر و دقیقتری برای فاصله بین دو نقطه مورد نظر باشد، می بایست علاوه بر تکرار مراحل اندازه گیری، مقدار حاصل از مرحله ۷ برای موارد مختلفی همچون خطای طول واقعی متر، انبساط طولی، انکسار و کمانی شدن متر نیز تصحیح گردد.

اندازه گیری فرایندی است توأم با تغییرات که ممکن است به دلایل مختلفی اتفاق بیافتند. به عنوان مثال چنانچه فاصله بین دو نقطه چندین بار بوسیله یک متر فولادی اندازه گیری شود و در همان حال تغییرات دمایی نیز وجود داشته باشد، متناظر با تغییرات دمایی شاهد تغییر در طول متر فولادی و در نتیجه تغییر در قرائت اعداد خواهیم بود. چنانچه بجز تغییر دما هیچ عامل دیگری وجود نداشته باشد، در آن صورت تغییرات موجود در اندازه گیری ها نشان دهنده تغییرات دمایی خواهند بود. اما واقعیت این است که عوامل بسیار زیادی در تغییرات اندازه گیری ها مؤثر هستند که دما یکی از آنهاست و عوامل دیگری نظیر نقص و محدودیتهای دستگاه های اندازه گیری و محدودیت توانایی عامل مشاهده برای سانتراژ، نشانه روی، تطابق و قرائت وجود دارند. بنابراین به دلیل تغییرات کوچکی که در هر مرحله از اندازه گیری اتفاق می افتد، نتیجه هیچ مشاهده ای نمی تواند عیناً تکرار شود.

از آنجا که تمام اندازه گیری ها توأم با تغییرات هستند، لذا کمیت مورد نظر بطور کامل و قطعی قابل تعیین نیست. حالت ایده آل آن است که برای کمیت مورد نظر به یک مقدار کاملاً ثابت و بدون تغییر دست یابیم که به آن "مقدار واقعی" گویند. اما در حقیقت آنچه ما بدست می آوریم چیزی جز یک "برآورد" از مقدار واقعی نیست. با توجه به اینکه تغییر در اندازه گیری ها یا مشاهدات دارای ماهیتی اتفاقی است و یک امری طبیعی تلقی می شود، از لحاظ ریاضی اندازه گیری یا مشاهده می بایست بعنوان یک متغیر در نظر گرفته شود. چنین متغیری در علم آمار و احتمال، متغیر تصادفی تصادفی نامیده می شود که در مورد آن بیشتر صحبت خواهد شد.

همانطور که اشاره شد حتی تحت شرایط محیطی کاملاً ثابت و یکسان باز هم تغییر در اندازه گیری ها دیده می شود، زیرا تغییر در مقادیر حاصل از اندازه گیری یک کمیت، پدیده ای کاملاً طبیعی محسوب می شود. بنابراین اختلافات بین مقادیر اندازه گیری شده با مقدار واقعی آن موسوم به "خطای اندازه گیری" نیز مورد انتظار است. ممکن است عبارت خطا القا کننده مفهوم اشتباه باشد که اینطور نیست و چنین خطایی به خطای بزرگ یا اشتباه معروف می باشد.

مطالعه خطاهای مشاهداتی و رفتار آنها اساساً معادل مطالعه خود مشاهدات است. به عبارت دیگر آنچه که بطور معمول از آن بعنوان "تئوری خطاها" یاد می شود معادل با "تئوری مشاهدات" است. اگر τ بیانگر مقدار واقعی یک کمیت (مانند طول یا زاویه) و x_i بیانگر یک اندازه گیری از آن باشد، آنگاه خطای اندازه گیری x_i بصورت زیر تعریف می شود.

$$\varepsilon_i = x_i - \tau \quad (1)$$

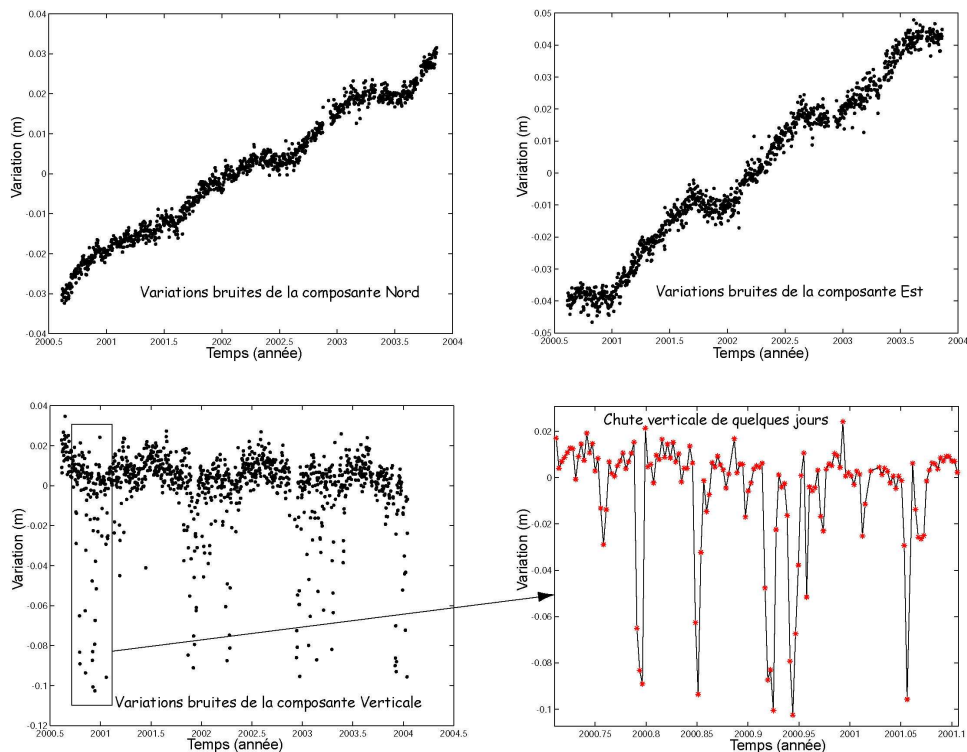
از آنجا که به طور معمول ما هرگز مقدار τ را نمی دانیم، لذا مقدار واقعی ε_i را نیز دست نیافتنی است. اما اگر بتوان برآورد خوبی از τ بدست آورد، آنگاه می توان از آن بعنوان یک معیار مناسب به منظور مطالعه تغییرات مشاهدات استفاده نمود. چنانچه برآورد τ را با $\hat{x}$ نمایش دهیم، آنگاه اختلاف بین برآورد $\hat{x}$ و مقدار مشاهده x_i را باقیمانده r_i نامیده و بصورت زیر تعریف می نمائیم.

$$r_i = x_i - \hat{x} \quad (2)$$

باقیمانده r_i کمیتی است که در عمل برای بیان و تجزیه و تحلیل تغییرات اندازه گیری ها مورد استفاده قرار می گیرد.

انواع خطاها

به طور سنتی می توان خطاهای اندازه گیری را به سه نوع دسته بندی نمود: خطاهای بزرگ، خطاهای سیستماتیک و خطاهای اتفاقی. هرچند انواع دیگری از دسته بندی خطاها وجود دارند، ولی این نوع دسته بندی در تئوری خطاها مناسب تر هستند. در شکل (۱) مثال خوبی از انواع رفتارهای سیستماتیک و اتفاقی در یک سری زمانی مولفه های مختصات ایستگاه GPS دائم تهران دیده می شود.



نگاره ۱- سری زمانی مولفه های مختصات ایستگاه GPS دائم تهران

• خطاهای بزرگ یا فاحش (Gross errors or Blunder)

خطاهای بزرگ نتیجه اشتباهاتی هستند که عمدتاً ناشی از بی توجهی عامل اندازه گیری می باشند. بعنوان مثال ممکن است عامل اندازه گیری اشتبهاً به هدف دیگری نشانه روی کند یا ممکن است در قرائت یا ثبت مشاهده خود دچار اشتباه شود، مثلاً ۴۱/۵۶ را بجای ۴۱/۶۵ ثبت کند. چنانچه عامل اندازه گیری فردی بی توجه و بی دقت باشد، بی شک اشتباهات زیادی در مشاهدات بوجود خواهند آمد. به هر حال در برآورد کمیت های مورد اندازه گیری، وجود اشتباهات غیر قابل قبول است و می بایست قبل از برآورد و تجزیه و تحلیل نتایج، اینگونه خطاهای بزرگ شناسایی و از لیست اندازه گیری ها حذف شوند.

معمولاً در پروژه های نقشه برداری که فقط جنبه سودمندی و اقتصادی آنها مورد توجه قرار گیرد، خطاهای بزرگ و اشتباهات نیز بیشتر اتفاق می افتند. از این رو روشهای صحرایی مناسبی به منظور کمک در کشف و تعیین اشتباهات وجود دارند که در زیر به برخی از آنها اشاره می شود.

✓ کنترل دقیق تمام قرائت ها روی نشانه های نقشه برداری

✓ تکرار قرائت ها و کنترل سازگاری معقول و منطقی بین آنها

- ✓ تصدیق اعداد ثبت شده با مقیاس اندازه گیری ها
 - ✓ تکرار کامل اندازه گیری ها بطور مستقل و کنترل آنها بمنظور وجود سازگاری بین آنها
 - ✓ استفاده از ابزار کنترل هندسی و جبری ساده مانند مجموع زاویای داخلی یک چند ضلعی
- اتخاذ روشهای فوق جهت جلوگیری از بروز اشتباهات بسیار مهم است ولی مجددا تاکید می شود چنانچه به هر دلیلی اشتباهی رخ دهد می بایست قبل از استفاده از اندازه گیری ها از لیست مشاهدات حذف شود.

• خطاهای سیستماتیک (Systematic errors)

خطاهایی که بر اساس روندهای قطعی و شناخته شده اتفاق افتاده و بوسیله مدل ها و روابط ریاضی قابل بیان باشند، خطاهای سیستماتیک نامیده می شوند. بعنوان مثال در مورد تغییر طول یک متر نواری فولادی با ضریب انبساط طولی مشخص، یک رابطه تابعی بین دما و تغییر طول نوار فولادی به صورت خطی ارائه شده است. بنابراین اگر طول متر نواری در یک دمای استاندارد معلوم باشد، تغییر طول نوار از حالت استاندارد، بدلیل انحراف از دمای استاندارد، بعنوان یک خطای سیستماتیک شناخته می شود.

ماهیت خطاهای سیستماتیک به گونه ای است که چنانچه اندازه گیری ها تحت شرایط کاملاً یکسانی تکرار شوند، خطاهای سیستماتیک نیز عیناً تکرار می شوند. به عنوان مثال در تکرار اندازه گیری یک فاصله با متر نواری فولادی، چنانچه از متر فولادی و عاملین اندازه گیری یکسانی تحت شرایط مشابه دما، کشش، شیب و..... استفاده شود، در آن صورت خطاهای سیستماتیک نیز عیناً تکرار خواهند شد.

اگر مقدار و علامت خطاهای سیستماتیک در سراسر فرآیند اندازه گیری ثابت باقی بماند به آن خطای "ثابت" اطلاق می شود. سازوکاری که رفتار خطاهای سیستماتیک تحت آن شکل می گیرد، ممکن است به عامل اندازه گیری، وسیله اندازه گیری، شرایط فیزیکی و محیطی در زمان اندازه گیری یا هر نوع ترکیبی از این عوامل وابسته باشد.

بعضی اوقات خطای شخصی عامل اندازه گیری نیز به سمت خطاهای سیستماتیک میل پیدا می کند. همانطور که خطای شخصی تحت شرایط مشاهده متغیر است، حواس طبیعی بینایی و شنوایی عامل اندازه گیری نیز ممکن است تغییر نماید.

ساختار غیر کامل دستگاه یا نقص در تنظیم دستگاه، خطای دستگاهی نامیده می شود که در زمره خطاهای سیستماتیک قرار می گیرد. ساختار غیر کامل شامل مواردی همچون عدم یکنواختی مقیاس درجه بندی و خروج از مرکزیت در سوار کردن قطعات دستگاه است. تنظیم دستگاهی ناقص نیز شامل مواردی همچون عمود نبودن محور تلسکوپ یک زاویه یاب بر محور افقی دستگاه است.

از آنجا که عموماً مشاهدات نقشه برداری در صحرا انجام می گیرند، متأثر از بسیاری عوامل فیزیکی و محیطی می باشند. برای مثال دما، کشش و شیب سطح زمین در متر کشی فواصل مؤثر هستند. در حالیکه علاوه بر دما، رطوبت و فشار هوا نیز در اندازه گیری فواصل با طولیاب های الکترونیکی، اندازه گیری زوایا و ترازایی مؤثر هستند. همه این اثرات به صورت تابعی از عوامل مذکور قابل بیان هستند و بنابراین بعنوان خطاهای سیستماتیک طبقه بندی می شوند.

تمام منابع خطاهای سیستماتیک که تا کنون مورد بحث قرار گرفتند، مستقیماً مرتبط با عمل اندازه گیری هستند. اما خطاهای سیستماتیک از طریق ساده سازی مدل های هندسی و ریاضی انتخابی نیز می توانند اتفاق بیفتند. بعنوان مثال اگر بجای یک مثلث کروی یک مثلث مسطحانی برای اتصال سه ایستگاه نقشه برداری با فواصل چند کیلومتری مورد استفاده قرار گیرد، اضافه کرویت بعنوان یک خطای سیستماتیک ظاهر خواهد شد.

با توجه به مطالب مطرح شده، باید گفت در پردازش و تجزیه و تحلیل اندازه گیری های نقشه برداری، باید تا حد امکان خطاهای سیستماتیک شناسایی و مورد تصحیح قرار گیرند. برای این منظور دو راه وجود دارد: بکارگیری روش های صحرائی و روش های محاسباتی. در روش های صحرائی با شناختی که از ماهیت هر خطای سیستماتیک وجود دارد روشی متناسب برای مقابله با آن اتخاذ می شود. به عنوان مثال می توان به قرار دادن تراز یاب در فاصله ای مساوی با شاخص اشاره کرد که باعث حذف خطاهای سیستماتیک انکسار، کرویت و کلیماسیون می شود. در روش های محاسباتی یا دفتری نیز با داشتن مدل یا رابطه ریاضی خطای

سیستماتیک مورد نظر، مانند خطاهای سیستماتیک ناشی از تغییرات جوی، کمائی شدن متر نواری یا شیب زمین، مقادیر تصحیح مربوط محاسبه و به اندازه گیری ها اعمال می شود.

• خطاهای اتفاقی (Random errors)

با فرض اینکه تمام اشتباهات شناسایی و حذف شوند و تصحیحات لازم نیز برای تمام خطاهای سیستماتیک شناخته شده اعمال گردند، هنوز برخی تغییرات در اندازه گیری ها دیده خواهند شد. این تغییرات باقیمانده ناشی از خطاهایی هستند که دیگر دارای یک رابطه مشخص ریاضی بر مبنای یک فرآیند قطعی نمی باشند. این خطاها دارای رفتاری کاملاً تصادفی هستند و لازم است بر همین اساس با آنها برخورد شود. قبلاً اشاره شد که یک اندازه گیری یا مشاهده از دیدگاه ریاضی یک متغیر است. از آنجا که هر اندازه گیری شامل مولفه های خطا با رفتار تصادفی است، لذا به وضوح می توان آن را یک "متغیر تصادفی" فرض نمود. بدیهی است که خطاهای تصادفی خودشان متغیرهای تصادفی هستند. البته یادآوری می شود در عمل بخشی از خطاهای سیستماتیک، که عموماً کوچک هستند و ناشناخته باقی می مانند، در قالب خطاهای اتفاقی ظاهر می شوند و بر همین اساس با آنها برخورد می شود که صحیح نیست. تنها راه مقابله با این نوع خطاها افزایش دقت عامل اندازه گیری، بکارگیری دستگاه های دقیقتر اندازه گیری و تکرار اندازه گیری ها می باشد که در کاهش سطح خطاهای اتفاقی موثر هستند. به خاطر داشته باشیم که با توجه به ماهیت این نوع خطاها، هیچگاه نمی توانیم آنها را به صفر برسانیم.

بر خلاف خطاهای سیستماتیک که در برخورد با آنها از مدل های ریاضی یا روابط تابعی استفاده می شود برای خطاهای تصادفی از مدل های احتمال استفاده می شود که در بخش ها و فصول آینده به بیان و شرح مختصری از مفاهیم اساسی تئوری احتمال پرداخته خواهد شد.

ارقام معنی دار و گرد کردن اعداد

در نمایش یک عدد، ارقام ۰ تا ۹ معنی دار محسوب می شوند. اما اگر رقم صفر برای نشان دادن محل ممیز باشد معنی دار محسوب نمی شود. مثلاً اعداد $۰/۰۹۳۴$ و $۰/۹۳۴$ هر دو دارای سه رقم معنی دار هستند و به این ترتیب عدد ۵۹۰۶ دارای چهار رقم معنی دار هستند.

یک مقدار عددی می بایست متشکل از تمام ارقام معین یا قطعی و اولین رقم نامعین یا مشکوک باشد. برای مثال اگر چهار رقم اول مقدار $۱۳۷/۸۲۴$ قطعی باشند و دو رقم آخر مشکوک باشند این مقدار عددی می بایست با پنج رقم معنی دار بیان شود یعنی: $۱۳۷/۸۲(۱,۳,۷)$ و ۸ قطعی هستند و ۲ مشکوک است).

تعیین تعداد ارقام معنی دار در کمیتی که مستقیماً اندازه گیری می شود معمولاً مشکل نیست، زیرا اساساً وابسته به کوچکترین درجه بندی دستگاه مورد استفاده است. مثلاً اگر یک فاصله با یک متر نواری با کوچکترین درجه بندی سانتیمتر و امکان تخمین تا یک میلیمتر، $۴۶۲/۵۱۳$ متر اندازه گیری شود، پنج رقم اول قطعی و رقم ششم به علت تخمینی بودن آن مشکوک است. بنابراین مقدار بدست آمده دارای شش رقم معنی دار است.

تعداد ارقام معنی دار در یک کمیت عددی با "گرد کردن" کاهش می یابد. چنانچه عمل گرد کردن بر اساس قواعد زیر انجام پذیرد کمترین خطا را در گرد کردن خواهیم داشت:

- ۱- اگر K رقم معنی دار مورد نیاز باشد، از تمام ارقام سمت راست رقم $K+1$ ام صرفنظر شود.
- ۲- رقم $K+1$ ام به شکل زیر مورد آزمایش قرار می گیرد.
 - الف- اگر بین صفر تا ۴ باشد از آن صرفنظر می شود، مثلاً $۱۲/۳۴۴۲۱$ در گرد کردن به چهار رقم معنی دار تبدیل به $۱۲/۳۴$ می شود.
 - ب- اگر بین ۶ تا ۹ است از آن صرفنظر نموده و رقم K ام یک واحد افزایش می یابد، مثلاً $۱/۳۷۶$ در گرد کردن به سه رقم بصورت $۱/۳۸$ در خواهد آمد.
 - ج- اگر ۵ است و رقم K ام زوج است از آن صرفنظر می شود، مثلاً گرد شده $۱۲/۳۴۵$ به چهار رقم $۱۲/۳۴$ است.

د- اگر ۵ است و رقم K ام فرد باشد از آن صرفنظر شده و یک واحد به رقم K ام اضافه می گردد، مثلاً ۱۲/۳۴۳۵ در گرد کردن به ۵ رقم معنی دار تبدیل به ۱۲/۳۴۴ می شود.

تعیین تعداد ارقام معنی دار یک کمیت حاصل از محاسبات در مقایسه با کمیتی که مستقیماً اندازه گیری می شود، خیلی ساده نیست. اگر چه تعدادی قواعد سودمند در این مورد وجود دارند، لیکن دو قاعده خیلی مهم زیر برای چهار عمل اصلی جمع، تفریق، ضرب و تقسیم بکار می روند.

در فرآیند عمل جمع (+) یا تفریق (-) می بایست حاصل عمل بگونه ای گرد شود که تعداد ارقام اعشاری آن برابر تعداد ارقام اعشاری عددی که دارای کمترین تعداد ارقام اعشاری است باشد. برای مثال مجموع زیر می بایست به ۳۸۲/۶ گرد شود زیرا ۱۴۹/۷ تنها دارای یک رقم بعد از ممیز است.

$$۱۶۵/۲۱ + ۱۴۹/۷ + ۶۵/۴۹۵ + ۲/۲۱۶۷ = ۳۸۲/۶۲۱۷$$

در فرآیند عمل ضرب (×) یا تقسیم (÷) تعداد ارقام معنی دار حاصل ضرب یا حاصل تقسیم می بایست برابر تعداد ارقام معنی دار عددی که دارای کمترین تعداد ارقام معنی دار است باشد. اعداد صحیح از این قاعده مستثنی هستند. برای مثال در هر دو حالت زیر تعداد ارقام معنی دار ۳ رقم هست زیرا تعداد ارقام معنی دار ۲/۱۵ سه رقم می باشد.

$$۲/۱۵ \times ۱۱/۱۲۳۴ = ۲۳/۹$$

$$۲ \times (۲/۱۵ \times ۱۱/۱۲۳۴) = ۴۷/۸ \quad (\text{عدد ۲ یک عدد صحیح است})$$

در محاسبات پیچیده و شلوغ که با تعداد خیلی زیادی از اعمال اصلی سروکار داریم، توصیه می شود از یک رقم معنی دار خیلی بزرگ در سراسر محاسبات استفاده شود و پس از اتمام محاسبات مطابق آنچه ارائه شد از عمل گرد کردن استفاده گردد.

فصل دوم

مروری بر مفاهیم اساسی آمار و احتمال

مفهوم احتمال در آزمایش‌هایی مطرح می‌شود که در نتایج آنها عدم قطعیتی وجود داشته باشد. یعنی، در مواردی که علیرغم تلاش در ثابت نگاه داشتن شرایط آزمایش، تغییر نتایج در تکرار آزمایش اجتناب ناپذیر است. اصطلاح آزمایش محدود به آزمایش‌هایی که در آزمایشگاه انجام می‌شوند نیست، بلکه شامل هر نوع فعالیتی است که منجر به گردآوری داده‌های مربوط به پدیده‌ای شود که مشاهدات تکراری، نتایج متفاوتی بدست می‌دهند. قلمرو کاربرد احتمال، تمام پدیده‌هایی را در برمی‌گیرد که در مورد آنها، نتایج را از پیش بطور دقیق نمی‌توان پیش بینی کرد. از آنجا که انواع مختلف مشاهدات نقشه برداری نیز در این دسته از آزمایشات قرار می‌گیرند، لذا آشنایی و بکارگیری آمار و احتمال نیز ضروری بنظر می‌رسد و قبل از طرح مباحث تخصصی مربوط به انتشار خطاهای مشاهدات نقشه برداری مروری بر مفاهیم اساسی آمار و احتمال خواهیم داشت.

برای روشن شدن مطلب، بحث خود را با مثالی از یک طول با تکرارهای زیاد اندازه‌گیری آغاز می‌کنیم که علاوه بر عاری بودن از خطاهای بزرگ، برای کلیه خطاهای سیستماتیک نیز تصحیح شده‌اند. بنابراین تغییرات باقیمانده در اندازه‌گیری‌ها تنها ناشی از خطاهای تصادفی هستند. اگر چه به هیچ وجه امکان تصحیح اندازه‌گیری‌ها برای خطاهای تصادفی وجود ندارد، لیکن امکان مطالعه رفتار آنها از طریق "توزیع فراوانی" میسر خواهد بود. به منظور تشکیل مدل احتمال اندازه‌گیری‌ها، توزیع فراوانی بعنوان یک اساس مورد استفاده قرار می‌گیرد.

توزیع فراوانی

یک طول حدوداً ۸۱۰ متری ۲۰۰ مرتبه اندازه‌گیری می‌شود. هیچ خطای بزرگی در اندازه‌گیری‌ها وجود نداشته و خطاهای سیستماتیک نیز تا دقت بهتر از یک سانتی متر تصحیح می‌شوند. نتیجه اندازه‌گیری‌های تصحیح شده که از ۸۱۰/۱۱ متر الی ۸۱۰/۲۳ متر نوسان دارند در جدول (۱) آمده است.

جدول ۱- نتیجه ۲۰۰ مرتبه اندازه گیری یک طول

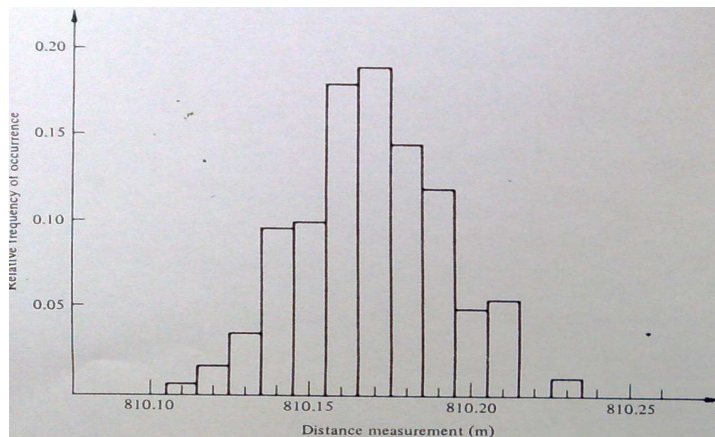
فراوانی	مقدار اندازه گیری (متر)
۱	۸۱۰/۱۱
۳	۸۱۰/۱۲
۷	۸۱۰/۱۳
۱۹	۸۱۰/۱۴
۲۰	۸۱۰/۱۵
۳۶	۸۱۰/۱۶
۳۸	۸۱۰/۱۷
۲۹	۸۱۰/۱۸
۲۴	۸۱۰/۱۹
۱۰	۸۱۰/۲۰
۱۱	۸۱۰/۲۱
۵	۸۱۰/۲۲
۲	۸۱۰/۲۳

به منظور ارزیابی و ترسیم فراوان نسبی برای مقادیر بدست آمده به شرح زیر عمل می کنیم. فراوانی نسبی با تقسیم تعداد تکرار هر مقدار بر کل تکرارهای اندازه گیری ها بدست می آید. با توجه به اینکه کل اندازه گیری ها ۲۰۰ بار می باشد، فراوانی های نسبی به صورت زیر در جدول (۲) بدست می آیند.

جدول ۲- فراوانی نسبی ۲۰۰ مرتبه اندازه گیری یک طول

فراوانی نسبی	مقدار اندازه گیری (متر)
۰/۰۰۵	۸۱۰/۱۱
۰/۰۱۵	۸۱۰/۱۲
۰/۰۳۵	۸۱۰/۱۳
۰/۰۹۵	۸۱۰/۱۴
۰/۱۰۰	۸۱۰/۱۵
۰/۱۸۰	۸۱۰/۱۶
۰/۱۹۰	۸۱۰/۱۷
۰/۱۴۵	۸۱۰/۱۸
۰/۱۲۰	۸۱۰/۱۹
۰/۰۵۰	۸۱۰/۲۰
۰/۰۵۵	۸۱۰/۲۱
۰/۰۲۵	۸۱۰/۲۲
۰/۰۱۰	۸۱۰/۲۳

بدیهی است با توجه به تعریف فراوانی نسبی بایستی مجموع فراوانی های نسبی برابر یک شود. فراوانی نسبی بدست آمده در نگاره (۱) بصورت پله ای (یا مستطیلی) ترسیم شده است.



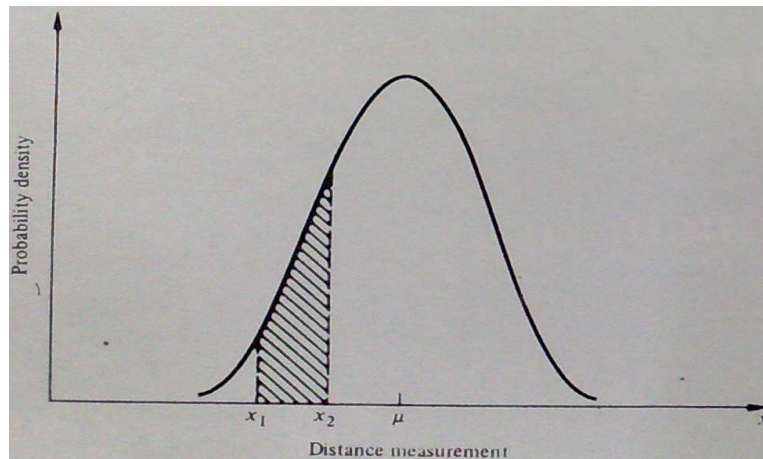
نگاره ۱- فراوانی نسبی اندازه گیری طول

نگاره (۱) بیانگر یک نمونه توزیع فراوانی در قالب یک هیستوگرام می باشد. عرض هر مستطیل بیانگر "فاصله دسته" و ارتفاع آن بیانگر فراوانی نسبی متناظر با آن می باشد. بعنوان مثال عرض بلند ترین مستطیل در نگاره (۱) بیانگر دسته ای از اندازه گیری ها است که بین 810.165 متر تا 810.175 متر واقع شده اند و ارتفاع آن نشان دهنده فراوانی نسبی متناظر با آن یعنی 0.190 است.

مطابق نگاره (۱) بیشترین فراوانی ها در حوالی مقدار 810.17 متر دیده می شود. چنانچه بتوان تعداد اندازه گیری ها را بی نهایت مرتبه افزایش داد، در اینصورت فراوانی های نسبی به یک حد پایدار نزدیک می شوند که این مقدار حدی معروف به "احتمال" است.

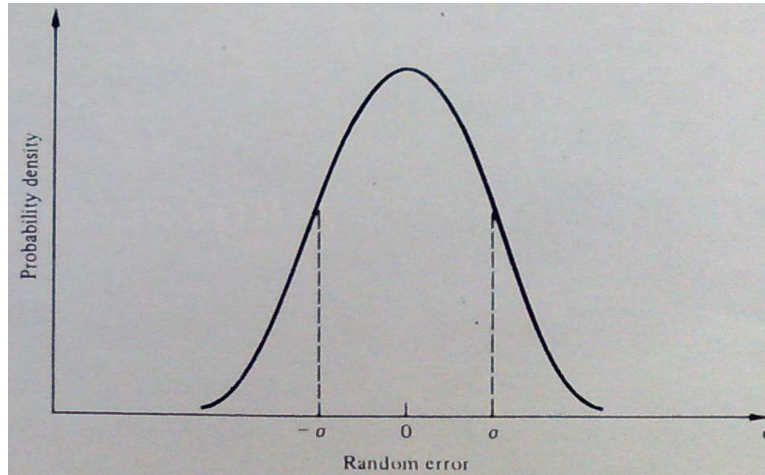
برای بیان احتمالات بجای استفاده از هیستوگرام اغلب از یک رابطه ریاضی، معروف به مدل احتمال یا توزیع احتمال، استفاده می شود. نمونه چنین مدلی در نگاره (۲) نمایش داده شده است. در این نگاره مقدار احتمال با مساحت زیر منحنی پیوسته به عنوان یک تابع ریاضی از اندازه گیری، بیان می گردد. بنابراین احتمال وقوع اندازه گیری بین دو مقدار x_1 و x_2 توسط مساحت هاشور خورده در نگاره (۲) نمایش داده می شود. همانطور که از قبل می دانیم، اندازه گیری یک متغیر تصادفی است و بنابراین منحنی مربوط به نگاره

(۲) به عنوان "تابع چگالی احتمال" متغیر تصادفی اندازه گیری شناخته می شود. مقدار مرکزی μ در نگاره (۲) بیانگر "مقدار متوسط" اندازه گیری است. اگر فرض شود هیچ خطای سیستماتیکی در اندازه گیری ها وجود نداشته باشد مقدار متوسط بعنوان مقدار واقعی در نظر گرفته می شود.



نگاره ۲- تابع چگالی احتمال اندازه گیری طول

یک تابع چگالی ساده مانند نگاره (۳) را می توان بعنوان مدل احتمال "خطای تصادفی" اندازه گیری نیز بکار برد. در این تابع چگالی مقدار متوسط برابر صفر است. مقدار متوسط μ به "موقعیت" یا "پارامتر موقعیت" توزیع احتمال معروف است. پارامتر مهم دیگر توزیع احتمال "انحراف معیار" (σ) است که مطابق نگاره (۳) بیانگر میزان پراکندگی توزیع احتمال است. بعنوان مثال اگر دو اندازه گیری A و B داشته باشیم و اندازه گیری A دارای تغییرات بیشتری نسبت به اندازه گیری B باشد، آنگاه اندازه گیری A انحراف معیار بزرگتری نسبت به اندازه گیری B خواهد داشت. مربع انحراف معیار σ معروف به "وریانس" σ^2 است.



نگاره ۳- تابع چگالی احتمال خطای تصادفی

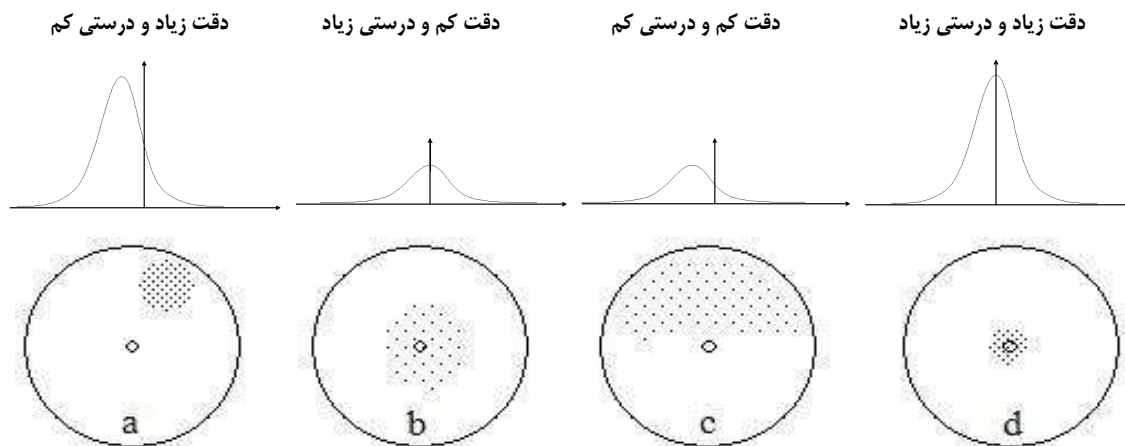
قابلیت اطمینان اندازه گیری ها

برای بیان قابلیت اطمینان اندازه گیری ها عبارات مختلفی بکار می رود که مشهور ترین آنها سه عبارت: دقت (precision)، درستی (accuracy) و فاصله اطمینان (confidence interval) هستند.

- دقت بیانگر درجه نزدیکی یا انطباق اندازه گیری های تکراری از یک کمیت واحد نسبت به همدیگر است. اگر اندازه گیری های تکراری با اختلافات بسیار جزئی در یک دسته قرار گیرند از دقت بالایی برخوردارند و اگر با اختلافات زیادی از یکدیگر دور باشند دارای دقت پایینی هستند. عموماً دقت بالا نتیجه توجه و مراقبت زیاد عامل اندازه گیری، انتخاب ابزار دقیق اندازه گیری و نیز بکارگیری روش های صحیح اندازه گیری می باشد. پارامتر دقت بوسیله پراکندگی توزیع احتمال قابل بیان می باشد. هر چه قدر شکل هندسی توزیع احتمال باریکتر باشد، بیانگر دقت بالاتر است. معیار متداول جهت نمایش دقت همان انحراف معیار (σ) است. بنابراین برای دقت های بالاتر مقدار σ کوچکتر و برای دقت های پایینتر مقدار σ بزرگتر است.

- درستی بیانگر درجه نزدیکی یا انطباق اندازه گیری ها با مقدار واقعی کمیت مورد اندازه گیری است. پارامتر درستی به نوعی بیانگر میزان حضور خطای سیستماتیک در اندازه گیری ها می باشد. بنابراین چنانچه هیچ اثر خطای سیستماتیکی وجود نداشته باشد، میانگین اندازه گیری ها می تواند بعنوان مقدار واقعی بکار رود.

- فاصله اطمینان دامنه یا بازه ای از اندازه گیری هاست که انتظار می رود اندازه واقعی در آن واقع شود. معمولاً برای بیان هر فاصله اطمینانی از یک سطح اطمینان خاص استفاده می شود. برای مثال فاصله اطمینان ۹۰٪، بازه ای است که احتمال وقوع اندازه واقعی در آن ۹۰ درصد است.
- تعبیر هندسی مفهوم واژه های دقت و درستی را می توان در رفتار تابع چگالی احتمال اندازه گیری و مطابقت آن با یک هدف تیراندازی در نگاره ۴ دید.



نگاره ۴- دقت و درستی در رفتار تابع چگالی احتمال اندازه گیری و یک هدف تیراندازی

فضای نمونه و پیشامد

مجموعه تمام نتایج متمایز ممکن در یک آزمایش تصادفی، فضای نمونه نامیده می شود و هر نتیجه متمایز یک پیشامد ساده یا یک عنصر فضای نمونه خوانده می شود. در اینجا فضای نمونه را با Ω نمایش می دهیم.

$$\Omega = \{u_1, u_2, u_3, \dots, u_n\}$$

بعنوان مثال می توان نتایج حاصل از ریختن یک تاس را به صورت زیر نمایش داد.

$$\Omega = \{1,2,3,4,5,6\}$$

فضاهای نمونه به دو دسته کلی زیر تقسیم می شوند.

- فضای نمونه گسسته: اگر فضای نمونه‌ای دارای تعداد عناصر متناهی یا بطور شمارش پذیر متناهی باشد را فضای نمونه گسسته گوییم.
- فضای نمونه پیوسته: اگر فضای نمونه شامل تمام اعداد متعلق به یک فاصله باشد آن را فضای نمونه پیوسته گوییم.

در یک فضای نمونه هر زیر مجموعه از آن را یک پیشامد، به مفهوم پیشامد تصادفی، می نامند. پیشامدی که به طور قطع رخ می دهد (پیشامد برابر با فضای نمونه) را پیشامد حتمی و پیشامدی که هیچ گاه رخ نمی دهد (پیشامد بدون عضو) را پیشامد ناممکن می نامند. به عنوان مثال در انداختن یک تاس پیشامد آمدن عدد ۷ یک پیشامد ناممکن و پیشامد آمدن عدد ۱ تا ۶ یک پیشامد حتمی است.

پیشامدها را دو به دو ناسازگار می گوئیم اگر تنها یکی از آنها به عنوان نتیجه آزمون بتواند رخ دهد. به عنوان مثال در بیرون آوردن یک مهره از یک ظرف محتوی مهره های قرمز و سیاه، بیرون آوردن مهره قرمز و بیرون آوردن مهره سیاه، ناسازگارند زیرا آنها به طور همزمان نمی توانند رخ دهند.

احتمال

احتمال وقوع یک پیشامد به معنای شانس وقوع آن پیشامد در انجام یک آزمایش تصادفی است. برای محاسبه احتمال عبارات مختلفی از جمله فراوانی نسبی و هم شانسی وجود دارد که هر کدام در برآورد احتمال می تواند مفید باشد. بنا بر تعریف تابع احتمال، تابعی از فضای نمونه به داخل مجموعه اعداد حقیقی بین ۰ تا ۱ است. به عبارت دقیقتر احتمال، تابعی است که به هر پیشامد از فضای نمونه عدد حقیقی $P(A)$ را به گونه ای نسبت می دهد که قوانین و شرایط زیر صدق کند. با فرض داشتن یک فضای نمونه n تایی احتمال وقوع

یک پیشامد عبارتست از نسبت پیشامدهای دلخواه (k پیشامد) به تمام پیشامدهای ممکن (n پیشامد) که رابطه آن بصورت زیر نمایش داده می شود:

$$P(A) = \frac{\sum_{i=1}^k (u_i)}{\sum_{i=1}^n (u_i)} \quad (۱)$$

قوانین و شرایط مدل احتمال

با توجه به رابطه فوق و توجه به اینکه تعداد پیشامدهای دلخواه همواره کوچکتر یا مساوی تعداد پیشامدهای ممکن است، می توان دریافت که احتمال وقوع یک پیشامد نیز همواره عددی بین صفر و یک است:

$$0 \leq P(A) \leq 1$$

$$P(\Omega) = 1 \quad \bullet \text{ احتمال وقوع یک پیشامد از تمام پیشامدهای ممکن:}$$

$$P(A \cap B) = P(A).P(B) \quad \bullet \text{ برای دو پیشامد ناسازگار (نابسته) A و B:}$$

$$P(A \cup B) = P(A) + P(B) - P(A \cap B) \quad \bullet \text{ اجتماع دو پیشامد A و B:}$$

$$P(A \cap B) = P(A) + P(B) - P(A \cup B) \quad \bullet \text{ اشتراک دو پیشامد A و B:}$$

$$P(A^c) = 1 - P(A) \quad \bullet \text{ متمم پیشامد A:}$$

$$P(A \cup A^c) = P(\Omega) = 1 = P(A) + P(A^c) \quad \bullet \text{ اجتماع پیشامد A و متمم آن:}$$

• اشتراک پیشامد A و متمم آن: $P(A \cap A^c) = P(\phi) = 0$

• احتمال پیشامد A به شرط پیشامد B : $P(A|B) = P(A \cap B)/P(B)$

متغیرهای تصادفی

اولین قدم در مدلسازی هر پدیده تصادفی تعریف یک متغیر تصادفی مناسب است. بنابر تعریف متغیر تصادفی تابعی است از فضای پیشامد ها به فضای اعداد حقیقی، یعنی به هر پیشامد یک عدد حقیقی نسبت داده می شود. با نسبت دادن اعداد به پیشامدها میتوان آنها را به صورت کمی بیان کرد، به عنوان مثال در آزمایش پرتاب سکه بجای ذکر نتیجه به صورت خط یا شیر می توانیم نتیجه را به صورت عددی بیان کنیم. متغیر تصادفی به دو صورت کلی گسسته و پیوسته تقسیم بندی می شود. متغیر تصادفی گسسته فقط مقادیر گسسته ای مانند ۰، ۱، ۲، ۳ را شامل می شود و معمولاً شمارش پذیر است. چنانچه مقادیر متغیر تصادفی، مجموعه ای از مقادیر حقیقی شمارش ناپذیر باشد آن را متغیر تصادفی پیوسته می نامند. برای تبیین بهتر موضوع، آزمایش ریختن همزمان دو تاس و پیشامدهای آن را در نظر می گیریم. متغیر تصادفی گسسته X را بعنوان جمع اعداد دو تاس در هر پیشامد در نظر می گیریم. بنابراین فضای نمونه جدید ما بصورت زیر تعریف می شود.

فضای نمونه: $\Omega_x = \{x_1, x_2, x_3, \dots, x_{11}\} = \{2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12\}$

فراوانی: $N: 1, 2, 3, 4, 5, 6, 5, 4, 3, 2, 1$

توابع توزیع احتمال گسسته

اگر X یک متغیر تصادفی گسسته باشد، تابعی که برای هر مقدار x در دامنه X $(f(x) = p(X = x))$ داده می شود، تابع توزیع احتمال X یا تابع چگالی احتمال X نامیده می شود. به عبارت دیگر توابع توزیع احتمالی که دارای متغیر تصادفی گسسته باشند، توابع توزیع گسسته نامیده می شوند.

چنانچه مجموع احتمالات کمتر از یک مقدار X در دامنه X ($F(x) = P(X \leq x)$) مورد نظر باشد آنگاه تابع توزیع چگالی تبدیل به تابع توزیع تجمعی یا تابع توزیع انباشته می شود. برای مثال آزمایش پرتاب دو تاس و متغیر تصادفی جمع اعداد دو تاس را در نظر بگیرید.

فضای نمونه: $\Omega_x = \{x_1, x_2, x_3, \dots, x_{11}\} = \{2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12\}$

فراوانی: $N: 1, 2, 3, 4, 5, 6, 5, 4, 3, 2, 1$

تابع چگالی احتمال: $f(x) = P(X = x): \frac{1}{36}, \frac{2}{36}, \dots, \frac{2}{36}, \frac{1}{36}$

تابع توزیع انباشته: $F(x) = P(X \leq x): \frac{1}{36}, \frac{3}{36}, \dots, \frac{35}{36}, \frac{36}{36}$

برخی از توابع توزیع گسسته که کاربرد و استفاده بیشتری در پدیده ها و زندگی روزمره ما دارند در زیر معرفی می شوند.

▪ **تابع توزیع یکنواخت:** ساده ترین توزیع احتمال گسسته است که متغیر تصادفی آن تمام مقادیرش دارای احتمال مساوی می باشد. بنابراین برای یک متغیر تصادفی X با n مقدار $(x_1, x_2, \dots, x_n)$ تابع توزیع یکنواخت بصورت زیر معرفی می گردد.

$$f(x) = \frac{1}{n} \quad (2)$$

▪ **تابع توزیع دو جمله ای:** اگر نتیجه یک آزمایش فقط با دو حالت موفقیت با احتمال p و عدم موفقیت با احتمال $q=1-p$ مشخص شود و متغیر تصادفی X در n بار آزمایش تعداد موفقیت ها در نظر گرفته شود، در آن صورت تابع توزیع احتمال دو جمله ای متغیر تصادفی X به صورت زیر خواهد بود.

$$f(x) = P(X = x) = C_n^x P^x q^{n-x} \quad (3)$$

$$x = 0, 1, 2, 3, \dots, n$$

همچنین تابع توزیع تجمعی، میانگین و واریانس متغیر تصادفی X برای $x = 0, 1, 2, 3, \dots, n$ بصورت زیر بدست آمده اند.

$$F(x) = P(X \leq x) = \sum_{i=0}^x C_n^i P^i q^{n-i} \quad \text{تابع توزیع تجمعی متغیر تصادفی } X:$$

$$\mu_x = \sum_{i=0}^x i \cdot C_n^i P^i q^{n-i} = np \quad \text{میانگین متغیر تصادفی } X:$$

$$\sigma_x^2 = \sum_{i=0}^x i^2 \cdot C_n^i P^i q^{n-i} - \mu_x^2 = npq \quad \text{واریانس متغیر تصادفی } X:$$

▪ **تابع توزیع برنولی:** در صورتیکه تعداد آزمایش های دوجمله ای برابر یک شود ($n=1$)، توزیع حاصل را برنولی می نامند. در واقع متغیر تصادفی X را که تنها دو مقدار صفر و یک را می پذیرد، متغیر برنولی و توزیع احتمال آن را توزیع برنولی می گویند.

▪ **تابع توزیع پواسن:** چنانچه متغیر تصادفی X بیانگر تعداد موفقیت اتفاقی در فاصله زمانی یا مکانی معینی باشد، تابع توزیع احتمال آن، معروف به تابع توزیع پواسن، به صورت زیر می باشد.

$$p(x, \mu) = f(x) = P(X = x) = \frac{e^{-\mu} \mu^x}{x!} \quad (4)$$

$$x = 0, 1, 2, 3, \dots, n$$

که در آن μ میانگین تعداد موفقیت اتفاقی و e عدد نمایی (≈ 2.71828) است. در صورتی که در توزیع دوجمله ای، تعداد آزمایش ها (n) خیلی بزرگ شود و احتمال موفقیت (p) خیلی کوچک شود، در آن صورت احتمال دوجمله ای و پواسن بسیار نزدیک به هم خواهند شد و در واقع توزیع پواسن تقریبی برای توزیع دوجمله ای خواهد بود.

توابع توزیع احتمال پیوسته

اگر X یک متغیر تصادفی پیوسته باشد، تابعی که برای هر مقدار x در دامنه X داده می‌شود، تابع توزیع احتمال X یا تابع چگالی احتمال X نامیده می‌شود. به عبارت دیگر توابع توزیع احتمالی که دارای متغیر تصادفی پیوسته باشند، توابع توزیع پیوسته نامیده می‌شوند. چنانچه مجموع احتمالات کمتر از یک مقدار x در دامنه X مورد نظر باشد آنگاه تابع توزیع چگالی تبدیل به تابع توزیع تجمعی یا تابع توزیع انباشته می‌شود. در واقع در بیان احتمال متغیرهای تصادفی گسسته از علامت جمع ($\sum$) و متغیرهای تصادفی پیوسته از علامت انتگرال ($\int$) استفاده می‌شود که می‌توان آن را به عنوان تفاوت اساسی بین حالات گسسته و پیوسته در نظر گرفت. اگر برای هر متغیر تصادفی پیوسته X ، تابع چگالی احتمال را با $f(x)$ (نگاره ۶) و تابع تجمعی احتمال را با $F(x)$ (نگاره ۵) نمایش دهیم و دو مقدار دلخواه a و b را با شرط $a \leq b$ در نظر بگیریم، در آن صورت روابط زیر برقرار خواهد بود.

$$f(x) \geq 0$$

$$f(x) = F'(x) = \frac{d}{dx} F(x)$$

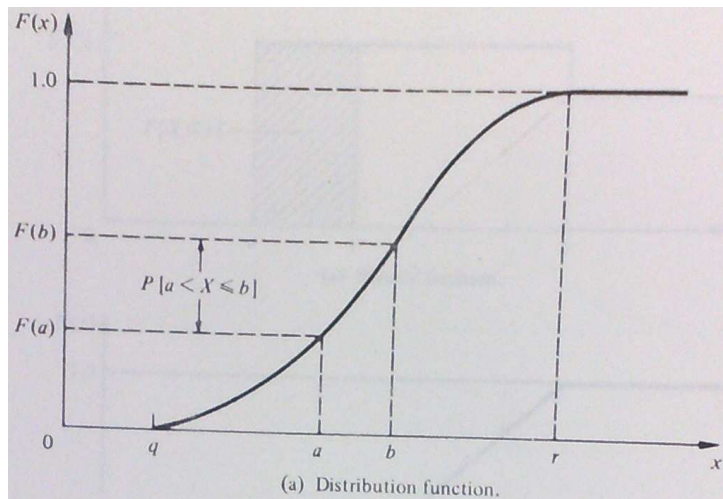
$$F(x) = P(X \leq x) = P(-\infty < X \leq x) = F(x) - F(-\infty) = \int_{-\infty}^x f(u) du$$

$$0 \leq F(x) \leq 1, F(a) \leq F(b)$$

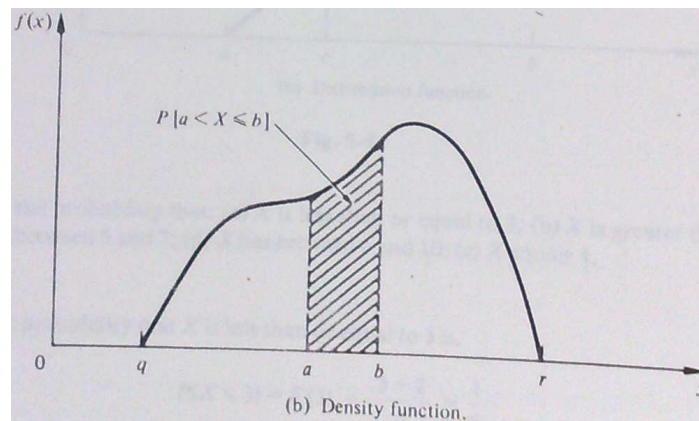
$$F(-\infty) = 0, F(\infty) = 1$$

$$\int_a^b f(x) dx = F(b) - F(a) = P(a \leq x \leq b)$$

$$\int_{-\infty}^{+\infty} f(x) dx = P(-\infty < x < \infty) = F(\infty) - F(-\infty) = 1$$



نگاره ۵- نمایش هندسی یک تابع تجمعی پیوسته



نگاره ۶- نمایش هندسی یک تابع چگالی پیوسته

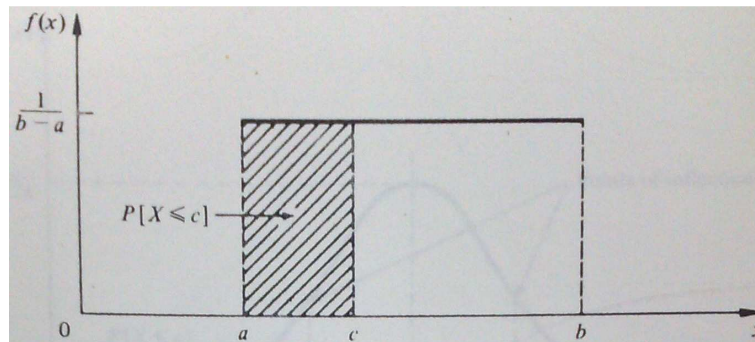
برخی از توابع توزیع پیوسته شناخته شده که استفاده بیشتری در موضوعات مهندسی و علوم دارند به ترتیب زیر معرفی می شوند.

▪ **تابع توزیع یکنواخت:** مطابق نگاره (۷)، ساده ترین متغیر تصادفی پیوسته که توزیع آن در

فاصله $[a, b]$ ثابت و در خارج از آن صفر باشد را دارای تابع توزیع یکنواخت می نامند و آن را به

صورت زیر معرفی می کنند.

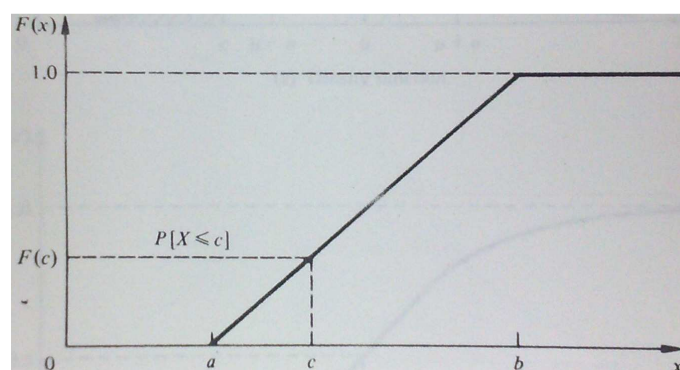
$$f(x) = \begin{cases} \frac{1}{b-a} & \text{for } a < x < b \\ 0 & \text{elsewhere} \end{cases} \quad (5)$$



نگاره ۷- نمایش هندسی تابع چگالی یکنواخت

بدیهی است با توجه به رابطه بین تابع چگالی و تابع تجمعی، تابع توزیع تجمعی یکنواخت نیز به صورت زیر بدست می آید (نگاره ۸).

$$F(x) = \begin{cases} 0 & \text{for } x \leq a \\ \int_a^x \frac{1}{b-a} du = \frac{x-a}{b-a} & \text{for } a < x < b \\ 1 & \text{for } x \geq b \end{cases} \quad (6)$$



نگاره ۷- نمایش هندسی تابع تجمعی یکنواخت

▪ **تابع توزیع نرمال:** توزیع نرمال، یکی از مهمترین توزیع ها در نظریه احتمال است و کاربردهای بسیاری در علم فیزیک و مهندسی دارد. این توزیع توسط کارل فریدریش گاوس در رابطه با کاربرد روش کمترین مربعات در آمارگیری کشف شد. فرمول آن بر حسب، دو پارامتر امید ریاضی و واریانس بیان می شود. همچنین تابع توزیع نرمال یا گاوس از مهمترین توابعی است که در مباحث آمار و احتمالات مورد بررسی قرار می گیرد چرا که به تجربه ثابت شده است که در دنیای اطراف ما توزیع بسیاری از متغیرهای طبیعی از همین تابع پیروی می کنند. بر همین اساس توزیع نرمال در مهندسی نقشه برداری نیز از جایگاه ویژه ای برخوردار است و رفتار کلیه اندازه گیری ها و مشاهدات نقشه برداری را به عنوان متغیرهای تصادفی نرمال X با این تابع ارزیابی می کنند. معادله ریاضی تابع چگالی و تجمعی احتمال نرمال با دو پارامتر میانگین (μ_x) و انحراف معیار (σ_x) متغیر تصادفی نرمال X به صورت زیر معرفی می شوند.

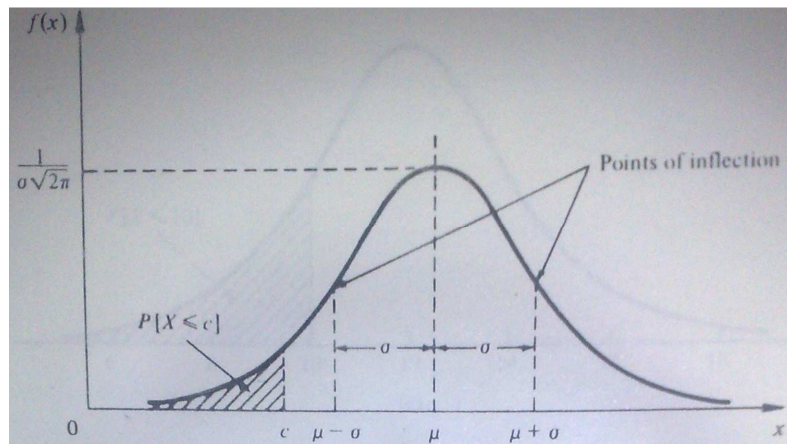
$$N(\mu_x, \sigma_x) = f(x) = \frac{1}{\sigma_x \sqrt{2\pi}} \exp\left(-\frac{(x - \mu_x)^2}{2\sigma_x^2}\right) \quad (7)$$

$$F(x) = \frac{1}{\sigma_x \sqrt{2\pi}} \int_{-\infty}^x \exp\left(-\frac{(u - \mu_x)^2}{2\sigma_x^2}\right) du \quad (8)$$

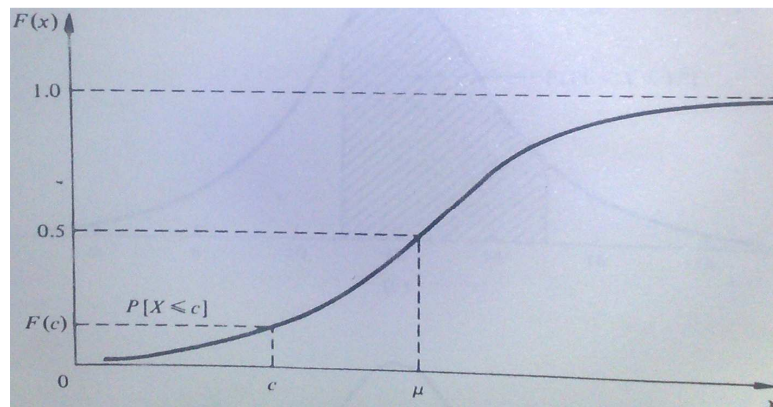
منحنی رفتار این تابع تا حد زیادی شبیه به زنگ های کلیسا می باشد و به همین دلیل به آن زنگوله ای هم گفته می شود. با وجود اینکه ممکن است ارتفاع و انحنای انواع مختلف این منحنی یکسان نباشد اما همه آنها یک ویژگی یکسان دارند و آن مساحت واحد زیر منحنی می باشد. این منحنی دارای خواص بسیار جالبی است از آن جمله که نسبت به محور عمودی متقارن می باشد، نیمی از مساحت زیر منحنی بالای مقدار میانگین μ_x و نیمه دیگر در پایین مقدار میانگین μ_x قرار دارد و اینکه هرچه از طرفین به میانگین نزدیک می شویم احتمال وقوع بیشتر می شود. فراگیری این تابع آنقدر زیاد است که دانشمندان اغلب برای مدل کردن متغیرهای تصادفی که با رفتار آنها آشنایی ندارند، از آن استفاده می کنند. بعنوان یک مثال در یک امتحان درسی نمرات دانش آموزان اغلب اطراف میانگین بیشتر می باشد و هر چه به سمت نمرات بالا یا پایین پیش برویم تعداد افرادی که این نمرات را گرفته اند کمتر می شود. این رفتار را

بسهولت می توان با یک توزیع نرمال مدل کرد. در نگاره های ۹ و ۱۰ تفسیر هندسی توابع چگالی و تجمعی احتمال نرمال به طور شماتیک نمایش داده شده است. برخی از ویژگی های مهم تابع چگالی نرمال به شرح زیر می باشد.

- ۱- سطح محصور بین منحنی تابع چگالی نرمال و محور طول ها برابر ۱ است.
- ۲- نقطه ماکزیمم منحنی در $X = \mu_x$ می باشد.
- ۳- نقاط عطف منحنی تابع چگالی نرمال است. $X = \pm\sigma_x$
- ۴- منحنی نسبت به خط عمودی $X = \mu_x$ متقارن است.
- ۵- در دو طرف حد متوسط μ_x ، منحنی به مجانب خود یعنی محور X ها نزدیک می گردد.



نگاره ۹- نمایش هندسی تابع چگالی نرمال با دو پارامتر (μ_x, σ_x)



نگاره ۱۰- نمایش هندسی تابع تجمعی نرمال با دو پارامتر (μ_x, σ_x)

چنانچه متغیر تصادفی X خود را به صورت زیر استاندارد نمائیم:

$$Z = \frac{X - \mu_x}{\sigma_x} \quad (9)$$

یعنی $\mu_z = 0$, $\sigma_z = 1$ در این صورت به آن تابع توزیع نرمال استاندارد گویند. در این حالت تابع توزیع چگالی $f(z)$ به صورت زیر خواهد بود.

$$f(z) = \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}} \quad (10)$$

تابع تجمعی نرمال استاندارد نیز، که بیشترین کاربرد را در مسایل مهندسی دارد، به صورت زیر می باشد.

$$F(z) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^z e^{-\frac{u^2}{2}} du \quad (11)$$

معمولاً مقادیر مشخصی از تابع تجمعی نرمال استاندارد محاسبه و در انتهای اغلب کتب مرتبط به صورت جداول می آیند. چنانچه مقدار احتمال تجمعی برای هر متغیر نرمال استاندارد z_α مورد نظر باشد، با توجه به روابط زیر براحتی از جداول مرتبط قابل استخراج است.

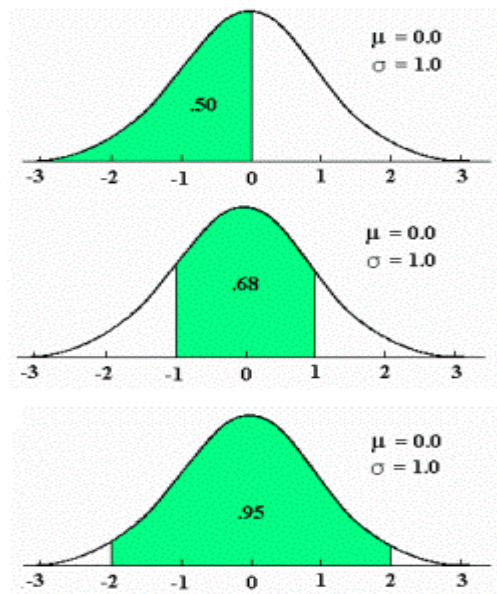
$$F(z_\alpha) = \int_{-\infty}^{z_\alpha} f(z) dz = P(Z \leq z_\alpha) = 1 - \alpha \quad (12)$$

$$1 - F(z_\alpha) = \int_{z_\alpha}^{+\infty} f(z) dz = P(Z \geq z_\alpha) = \alpha$$

که در آن α سطح زیر منحنی از نقطه z_α تا $+\infty$ است. بدیهی است چنانچه هدف محاسبه احتمال بین دو مقدار z_{α_1} و z_{α_2} باشد، با استفاده از جداول مذکور به سادگی قابل محاسبه است.

$$P(z_{\alpha_1} \leq Z \leq z_{\alpha_2}) = F(z_{\alpha_2}) - F(z_{\alpha_1}) \quad (۱۳)$$

در زیر چند مثال مربوط به سطح زیر منحنی نرمال استاندارد که از اهمیت بیشتری در تجزیه و تحلیل مسایل آماری برخوردارند، آمده اند.



نگاره ۱۱- سطح زیر منحنی نرمال استاندارد برای سه حالت متفاوت

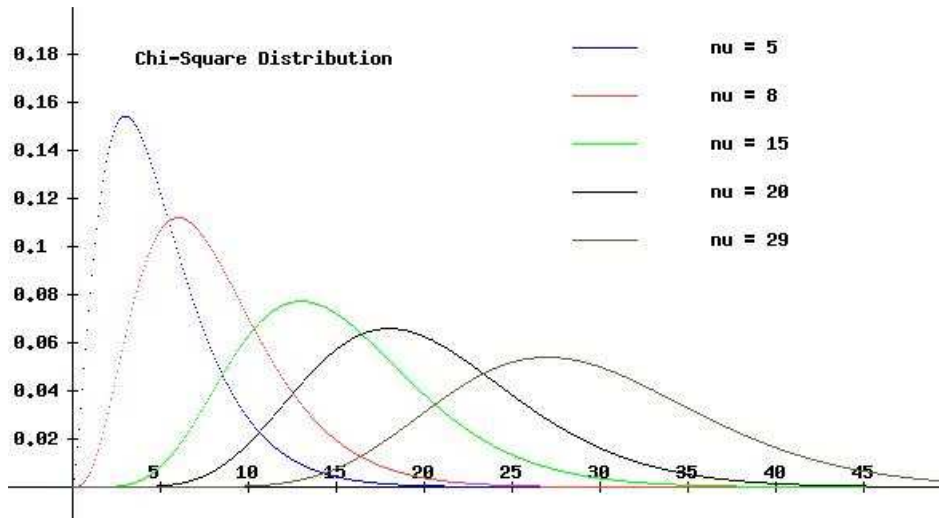
▪ **تابع توزیع کای اسکور (χ^2):** چنانچه n متغیر تصادفی نرمال استاندارد $(z_1, z_2, z_3, \dots, z_n)$

داشته باشیم و مجموع مربعات آنها را به عنوان یک متغیر تصادفی جدید کای اسکور

($\chi_n^2 = z_1^2 + z_2^2 + \dots + z_n^2$) در نظر بگیریم، در آن صورت توزیع احتمال این متغیر تصادفی جدید را

تابع توزیع کای اسکور با درجه آزادی n ($f(\chi_n^2)$) می نامند. در نگاره ۱۲ تفسیر هندسی تابع چگالی

کای اسکور برای درجات مختلف آزادی نمایش داده شده است.



نگاره ۱۲- نمایش هندسی تابع چگالی کای اسکور برای درجات مختلف آزادی

مقدار احتمال برای مقدار کوچکتر از $\chi^2_{n;\alpha}$ (تابع تجمعی کای اسکور)، میانگین و وریانس متغیر تصادفی کای اسکور با درجه آزادی n (χ^2_n) نیز به صورت زیر ارائه می شوند.

$$F(\chi^2_n) = \int_{-\infty}^{\chi^2_{n;\alpha}} f(\chi^2_n) d\chi^2_n = P(\chi^2_n \leq \chi^2_{n;\alpha}) = 1 - \alpha \quad (۱۴)$$

$$\mu_{\chi^2_n} = n$$

$$\sigma_{\chi^2_n}^2 = 2n$$

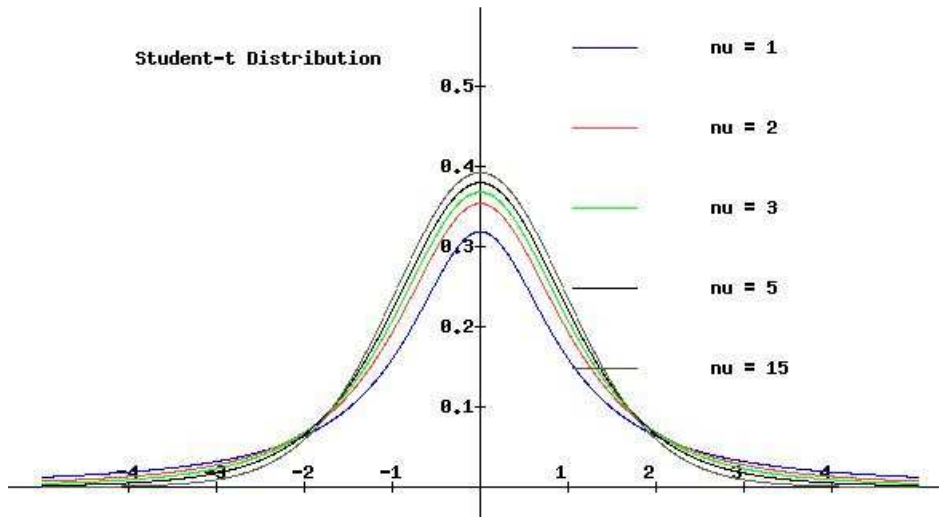
که در آن α سطح زیر منحنی از نقطه $\chi^2_{n;\alpha}$ تا $+\infty$ است. مقادیر احتمال کای اسکور نیز مانند توزیع نرمال در جداول مخصوصی تهیه و انتهای کتب مرتبط گنجانده شده اند.

▪ **توزیع احتمال t-student:** چنانچه ترکیبی از متغیرهای تصادفی نرمال استاندارد Z و کای

اسکور χ^2_n با درجه آزادی n بصورت $t_n = Z / \sqrt{\chi^2_n / n}$ در نظر گرفته شود، در این صورت متغیر

تصادفی جدید t_n را متغیر تصادفی t-student با درجه آزادی n و توزیع احتمال آن را نیز تابع

چگالی t-student با درجه آزادی n می نامند. در نگاره ۱۳ تفسیر هندسی تابع چگالی t-student برای درجات مختلف آزادی نمایش داده شده است.



نگاره ۱۳- نمایش هندسی تابع چگالی t-student برای درجات مختلف آزادی

مقدار احتمال برای مقدار کوچکتر از $t_{n;\alpha}$ (تابع تجمعی t-student)، میانگین و واریانس متغیر تصادفی t-student با درجه آزادی n (t_n) نیز به صورت زیر ارائه می شوند.

$$F(t_{n;\alpha}) = \int_{-\infty}^{t_{n;\alpha}} f(t_n) dt_n = P(t_n \leq t_{n;\alpha}) = 1 - \alpha \quad (15)$$

$$\mu_{t_n} = 0 \quad \text{for } n > 1$$

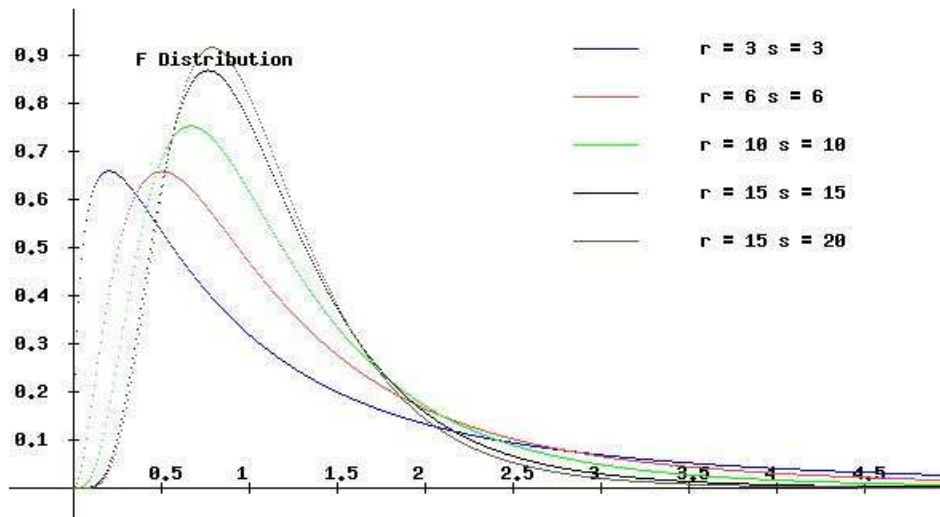
$$\sigma_{t_n}^2 = \frac{n}{n-2} \quad \text{for } n > 2$$

که در آن α سطح زیر منحنی از نقطه $t_{n;\alpha}$ تا $+\infty$ است. مقادیر احتمال t-student نیز مانند توابع نرمال و کای اسکور در جداول مخصوصی تهیه و انتهای کتب مرتبط گنجانده شده اند.

▪ **توزیع احتمال F:** چنانچه ترکیبی از دو متغیر تصادفی کای اسکور χ_{n1}^2 و χ_{n2}^2 با درجات آزادی

$n1$ و $n2$ بصورت $F_{n1, n2} = \chi_{n1}^2 / \chi_{n2}^2$ در نظر گرفته شود، در این صورت متغیر تصادفی

جدید F_{n_1, n_2} را متغیر تصادفی F با درجه آزادی n_1 و n_2 و توزیع احتمال آن را نیز تابع چگالی F (فیشر) با درجه آزادی n_1 و n_2 می نامند. در نگاره ۱۴ تفسیر هندسی تابع چگالی F برای درجات مختلف آزادی نمایش داده شده است.



نگاره ۱۴- نمایش هندسی تابع چگالی F برای درجات مختلف آزادی

مقدار احتمال برای مقدار کوچکتر از $F_{n_1, n_2; \alpha}$ (تابع تجمعی F)، میانگین و وریانس متغیر تصادفی F با درجه آزادی n_1 و n_2 (F_{n_1, n_2}) نیز به صورت زیر ارائه می شوند.

$$F(F_{n_1, n_2; \alpha}) = \int_{-\infty}^{F_{n_1, n_2; \alpha}} f(F_{n_1, n_2}) dF_{n_1, n_2} = P(F_{n_1, n_2} \leq F_{n_1, n_2; \alpha}) = 1 - \alpha \quad (16)$$

$$\mu_{F_{n_1, n_2}} = \frac{n_2}{n_2 - 2} \quad \text{for } n_2 > 2$$

$$\sigma_{F_{n_1, n_2}}^2 = \frac{2n_2^2(n_1 + n_2 - 2)}{n_1(n_2 - 2)^2(n_2 - 4)} \quad \text{for } n_2 > 4$$

که در آن α سطح زیر منحنی از نقطه $F_{n_1, n_2; \alpha}$ تا $+\infty$ است. مقادیر احتمال F نیز مانند توابع نرمال، کای اسکور و t-student در جداول مخصوصی تهیه و انتهای کتب مرتبط گنجانده شده اند.

توزیع احتمال چند متغیره

در صورتی که یک بردار شامل چند متغیر تصادفی مانند $\underline{X} = (x, y, z, \dots, t)$ ، تابع چگالی توام آنها با

$$f(\underline{X}) = f(x, y, z, \dots, t)$$

با مفهوم دیفرانسیلی به صورت زیر تعریف می شود.

$$P(\underline{X}_1 \leq \underline{X} \leq \underline{X}_1 + d\underline{X}) = f(\underline{X})d\underline{X}$$

$$P(x_1 \leq X \leq x_1 + dx, y_1 \leq Y \leq y_1 + dy, z_1 \leq Z \leq z_1 + dz, \dots, t_1 \leq T \leq t_1 + dt) = f(x_1, y_1, z_1, \dots, t_1) dx dy dz \dots dt$$

همچنین تابع تجمعی احتمال چند متغیره $F(\underline{X}_1) = F(x_1, y_1, z_1, \dots, t_1)$ نیز به صورت زیر بدست می آید.

$$F(\underline{X}_1) = F(x_1, y_1, z_1, \dots, t_1) = \int_{-\infty}^{x_1} \int_{-\infty}^{y_1} \int_{-\infty}^{z_1} \dots \int_{-\infty}^{t_1} f(x_1, y_1, z_1, \dots, t_1) dx dy dz \dots dt$$

برای سادگی مطلب حالت دو بعدی با دو متغیر تصادفی (X, Y) در نظر گرفته و ارتباط بین دو تابع چگالی و تجمعی توام را یادآوری می کنیم.

$$f(x, y) = \frac{\partial^2}{\partial x \partial y} F(x, y)$$

$$F(x, y) = \int_{-\infty}^x \int_{-\infty}^y f(u, v) du dv$$

یکی از مواردی که در توابع چند متغیره مطرح و مورد علاقه می باشد، توابع کناری یا حاشیه ای است که برای حالت دو بعدی به صورت زیر قابل استخراج هستند.

$$f(x) = \int_{-\infty}^{\infty} f(x, y) dy$$

$$f(y) = \int_{-\infty}^{\infty} f(x, y) dx$$

در صورت وجود استقلال بین متغیرهای تصادفی در حالت چند بعدی، توابع چگالی و تجمعی احتمال آنها به شکل بسیار ساده تری با ضرب توابع چگالی و تجمعی تک تک متغیرهای تصادفی جایگزین می شوند.

$$f(x, y, z, \dots, t) = f(x) \cdot f(y) \cdot f(z) \dots f(t)$$

$$F(x, y, z, \dots, t) = F(x) \cdot F(y) \cdot F(z) \dots F(t)$$

وقتی که صحبت از بررسی همزمان چند متغیر تصادفی می کنیم، بحث وابستگی و عدم وابستگی بین دو به دوی این متغیرها نیز مطرح می شود که عموماً بوسیله کوریانس یا ضریب همبستگی بیان می شود. از آنجا که پرداختن به این موضوع وابسته به تعریف امید ریاضی است، لذا ابتدا به امید ریاضی پرداخته و سپس به دنبال مفهوم وابستگی می رویم.

▪ **تابع چگالی نرمال چند متغیره:** با فرض n متغیر تصادفی $(X_1, X_2, \dots, X_n)$ ، تابع چگالی توام نرمال آنها به صورت زیر معرفی می شود.

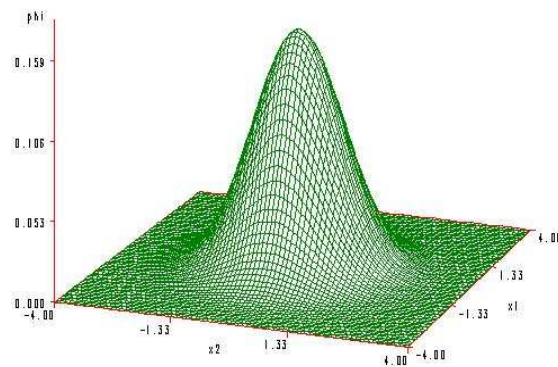
$$f(x_1, x_2, \dots, x_n) = \frac{1}{(2\pi)^{n/2} \cdot |C_x|^{1/2}} \times \exp\left[-\frac{1}{2}(\underline{X} - \underline{\mu}_x)^T \underline{C}_x^{-1}(\underline{X} - \underline{\mu}_x)\right] \quad (17)$$

چنانچه تنها دو متغیر تصادفی X و Y در نظر گرفته شوند، تابع چگالی توام نرمال آنها به صورت زیر خواهد بود.

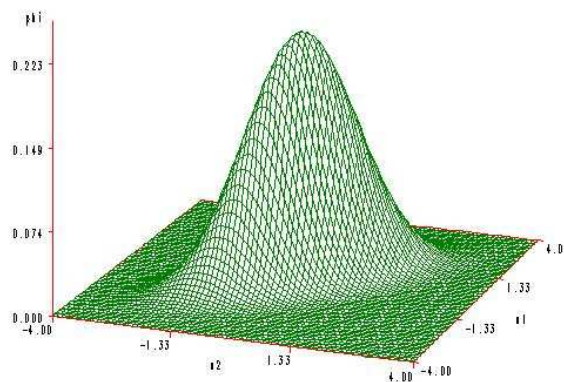
$$f(x_1, x_2) = \frac{1}{2\pi\sqrt{\sigma_1^2\sigma_2^2 - \sigma_{12}^2}} \exp\left[\frac{-\sigma_1^2\sigma_2^2}{2(\sigma_1^2\sigma_2^2 - \sigma_{12}^2)} \times \left(\left(\frac{x_1 - \mu_1}{\sigma_1}\right)^2 - 2\sigma_{12} \frac{(x_1 - \mu_1)(x_2 - \mu_2)}{\sigma_1^2\sigma_2^2} + \left(\frac{x_2 - \mu_2}{\sigma_2}\right)^2\right)\right] \quad (18)$$

در نگاره (۱۵) سه نمایش مختلف از تابع چگالی نرمال دو متغیره برای ضریب همبستگی ۰، ۰.۷ و ۰.۹ آمده است.

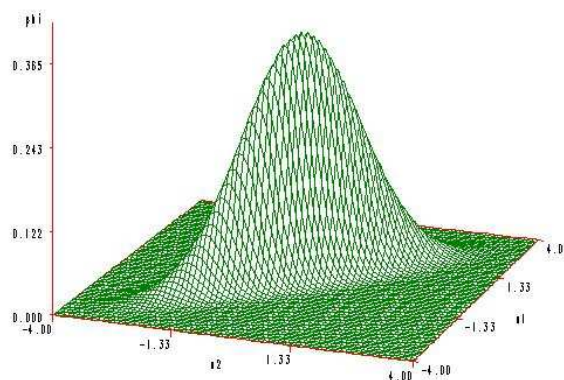
Bivariate Normal Density — $r=0.0$



Bivariate Normal Density — $r=0.7$



Bivariate Normal Density — $r=0.9$



نگاره ۱۵- نمایش هندسی تابع چگالی نرمال دو متغیره برای ضریب همبستگی ۰، ۰.۷ و ۰.۹

امید ریاضی

بنا بر تعریف امید ریاضی هر تابعی مانند $g(X)$ که در آن X متغیر تصادفی و دارای تابع چگالی $f(x)$ است، بصورت زیر تعریف می شود.

$$E[g(X)] = \int_{-\infty}^{\infty} g(x) f(x) dx \quad (19)$$

برخی قواعد در استفاده از عملگر امید ریاضی به صورت زیر معرفی می شوند.

$$E(X + Y) = E(X) + E(Y)$$

$$E(cX) = cE(X)$$

$$E(XY) = E(X)E(Y) + E[(X - \mu_x)(Y - \mu_y)]$$

در صورتیکه متغیرهای تصادفی X و Y نایسته باشند، آنگاه خواهیم داشت.

$$E(XY) = E(X)E(Y)$$

با استفاده از تعریف امید ریاضی، ممان مرتبه i ام متغیر تصادفی X با در نظر گرفتن $g(X) = X^i$ ، بصورت زیر تعریف می شود.

$$E[X^i] = \int_{-\infty}^{\infty} x^i f(x) dx \quad (20)$$

معروفترین ممان با تعریف فوق، ممان اول متغیر تصادفی X است که میانگین تئوریک یا واقعی آن متغیر تصادفی نامیده می شود ($E[X] = \mu_x$). نوع دیگری از ممان، ممان مرکزی نامیده می شود که برای یک متغیر تصادفی X با در نظر گرفتن $g(X) = (X - \mu_x)^i$ ، بصورت زیر بدست می آید.

$$E[(X - \mu_x)^i] = \int_{-\infty}^{\infty} (x - \mu_x)^i f(x) dx \quad (21)$$

چند ممان مرکزی نخست بصورت زیر تعریف می شوند که از بین آنها ممان مرکزی دوم بعنوان وریانس تئوریک یا واقعی از اهمیت بیشتری در آمار و احتمال برخوردار است.

$$\begin{aligned} C_1 &= E[(X - \mu_x)] = 0 \\ C_2 &= E[(X - \mu_x)^2] = E[X^2] + \mu_x^2 - 2E(\mu_x X) = E[X^2] - \mu_x^2 = \sigma_x^2 \\ C_3 &= E[(X - \mu_x)^3] = E[X^3] - 3\mu_x E[X^2] + 2\mu_x^2 \end{aligned} \quad (22)$$

در حالت گسسته که با متغیرهای تصادفی گسسته سروکار داریم بجای علامت انتگرال ($\int$) از علامت جمع ($\sum$) استفاده می کنیم.

کورینانس و ضریب همبستگی

اکنون که با مفهوم و برخی قواعد امید ریاضی آشنا شدیم، می توانیم کورینانس بین دو متغیر تصادفی X و Y ($Cov(X, Y) = \sigma_{xy}$) را به صورت زیر تعریف نماییم.

$$\begin{aligned} \sigma_{xy} &= E[(X - \mu_x)(Y - \mu_y)] \\ &= E[XY - \mu_x Y - \mu_y X + \mu_x \mu_y] \\ &= E[XY] - \mu_x E[Y] - \mu_y E[X] + \mu_x \mu_y \\ &= E[XY] - \mu_x \mu_y \end{aligned} \quad (23)$$

مقدار عددی کوریانس دو متغیر تصادفی نمی تواند به روشنی بیانگر میزان وابستگی بین آنها باشد. بر همین اساس امکان مقایسه میزان وابستگی بین زوج متغیرهای تصادفی وجود نیز وجود ندارد. بنابراین یک کمیت جدید بنام ضریب همبستگی بین دو متغیر تصادفی X و Y که دارای مفهوم نسبی است معرفی و به جای کوریانس استفاده می شود. رابطه ضریب همبستگی به صورت زیر می باشد و مقدار عددی آن در فاصله $[-1,1]$ قرار دارد.

$$\rho_{xy} = \frac{\sigma_{xy}}{\sigma_x \cdot \sigma_y} \quad (24)$$

که در آن σ_{xy} کوریانس بین دو متغیر تصادفی X و Y ، σ_x انحراف معیار متغیر تصادفی X و σ_y انحراف معیار متغیر تصادفی Y است.

امید ریاضی توابع

استفاده از امید ریاضی و توسعه آن برای برخی مدل های ریاضی ما را در درک مفهوم انتشار خطاها بسیار یاری می رساند. در زیر هر دو حالت خطی و غیرخطی را برای چنین مدل هایی مورد بررسی قرار می دهیم.

▪ **مدل خطی:** برای سادگی و درک بهتر مطلب یک مدل خطی که تابعی از متغیرهای تصادفی X و Y است را در نظر می گیریم و سپس امید ریاضی را بر روی آن اعمال می کنیم.

$$Z = aX + bY$$

$$\sigma_z^2 = E[(Z - \mu_z)^2] = a^2 \sigma_x^2 + b^2 \sigma_y^2 + 2ab \sigma_{xy}$$

در صورتیکه متغیرهای تصادفی X و Y ناپسته باشند، در آن صورت خواهیم داشت.

$$\sigma_z^2 = a^2 \sigma_x^2 + b^2 \sigma_y^2$$

چنانچه حالت ساده فوق را به یک تابع n متغیری بسط دهیم و فرض استقلال بین آنها برقرار باشد، آنگاه خواهیم داشت.

$$Y = a_1 X_1 + a_2 X_2 + a_3 X_3 + \dots + a_n X_n \quad (25)$$

$$\sigma_y^2 = a_1^2 \sigma_{x1}^2 + a_2^2 \sigma_{x2}^2 + a_3^2 \sigma_{x3}^2 + \dots + a_n^2 \sigma_{xn}^2$$

با توجه به روابط فوق در می یابیم که چگونه با استفاده از امید ریاضی می توان اثر خطاهای اندازه گیری را در یک مدل خطی روی نتیجه نهایی دید.

▪ **مدل غیر خطی:** برای شروع بررسی یک مدل غیرخطی که تابعی از سه متغیر تصادفی Y, X و Z با میانگین های μ_x, μ_y, μ_z و واریانس های $\sigma_x^2, \sigma_y^2, \sigma_z^2$ است را به صورت زیر در نظر می گیریم.

$$U = f(x, y, z)$$

اگر تابع را حول مقادیر میانگین μ_x, μ_y, μ_z بسط تیلور دهیم، می توانیم با استفاده از امید ریاضی به واریانس تابع برسیم.

$$U = f(\mu_x, \mu_y, \mu_z) + \frac{\partial f}{\partial x}(X - \mu_x) + \frac{\partial f}{\partial y}(Y - \mu_y) + \frac{\partial f}{\partial z}(Z - \mu_z) + \dots$$

$$U - E(U) = \frac{\partial f}{\partial x}(X - \mu_x) + \frac{\partial f}{\partial y}(Y - \mu_y) + \frac{\partial f}{\partial z}(Z - \mu_z) + \dots$$

$$\sigma_u^2 = E[(U - E(U))^2] = \left(\frac{\partial f}{\partial x}\right)^2 E[(X - \mu_x)^2] + 2 \frac{\partial f}{\partial x} \cdot \frac{\partial f}{\partial y} E[(X - \mu_x)(Y - \mu_y)] + 2 \frac{\partial f}{\partial x} \cdot \frac{\partial f}{\partial z} E[(X - \mu_x)(Z - \mu_z)] + \left(\frac{\partial f}{\partial y}\right)^2 E[(Y - \mu_y)^2] + 2 \frac{\partial f}{\partial y} \cdot \frac{\partial f}{\partial z} E[(Y - \mu_y)(Z - \mu_z)] + \left(\frac{\partial f}{\partial z}\right)^2 E[(Z - \mu_z)^2] + \dots$$

$$\sigma_u^2 = \left(\frac{\partial f}{\partial x}\right)^2 \sigma_x^2 + 2 \frac{\partial f}{\partial x} \cdot \frac{\partial f}{\partial y} \sigma_{xy} + 2 \frac{\partial f}{\partial x} \cdot \frac{\partial f}{\partial z} \sigma_{xz} + \left(\frac{\partial f}{\partial y}\right)^2 \sigma_y^2 + 2 \frac{\partial f}{\partial y} \cdot \frac{\partial f}{\partial z} \sigma_{yz} + \left(\frac{\partial f}{\partial z}\right)^2 \sigma_z^2 + \dots$$

اگر متغیرهای تصادفی X , Y و Z نایسته باشند که در اغلب اندازه گیری های ما این فرض وجود دارد به رابطه ساده تر زیر می رسیم.

$$\sigma_u^2 \approx \left(\frac{\partial f}{\partial x}\right)^2 \sigma_x^2 + \left(\frac{\partial f}{\partial y}\right)^2 \sigma_y^2 + \left(\frac{\partial f}{\partial z}\right)^2 \sigma_z^2$$

با فرض یک تابع غیرخطی n متغیری و نایسته بودن آنها، می توان انتشار خطاهای اتفاقی متغیرهای تصادفی را بر روی نتیجه نهایی به صورت زیر معرفی کرد.

$$Y = f(X_1, X_2, X_3, \dots, X_n) \quad (26)$$

$$\sigma_y^2 = \left(\frac{\partial f}{\partial x_1}\right)^2 \sigma_{x1}^2 + \left(\frac{\partial f}{\partial x_2}\right)^2 \sigma_{x2}^2 + \left(\frac{\partial f}{\partial x_3}\right)^2 \sigma_{x3}^2 + \dots + \left(\frac{\partial f}{\partial x_n}\right)^2 \sigma_{xn}^2$$

ماتریس های وریانس-کوریانس

یک بردار متغیر تصادفی $\underline{X}$ و اختلاف آن با بردار میانگین واقعی اش $(\underline{X} - \underline{\mu})$ به صورت بردارهای ستونی زیر قابل نمایش هستند.

$$\underline{X} = \begin{bmatrix} x_1 \\ x_2 \\ x_3 \\ \vdots \\ x_n \end{bmatrix}, \quad (\underline{X} - \underline{\mu}) = \begin{bmatrix} (x_1 - \mu_1) \\ (x_2 - \mu_2) \\ (x_3 - \mu_3) \\ \vdots \\ (x_n - \mu_n) \end{bmatrix}$$

از آنجا که توان دوم عناصر بردار $(\underline{X} - \underline{\mu})$ به صورت یک عنصر اسکالر قابل انجام و نمایش نیست، لذا در جبر ماتریسی فرم ماتریسی حاصلضرب $(\underline{X} - \underline{\mu})(\underline{X} - \underline{\mu})^T$ برای آن پیشنهاد شده است که حاصل آن به صورت زیر می باشد.

$$(\underline{X} - \underline{\mu})(\underline{X} - \underline{\mu})^T = \begin{bmatrix} (x_1 - \mu_1)^2 & (x_1 - \mu_1)(x_2 - \mu_2) & (x_1 - \mu_1)(x_3 - \mu_3) & \dots & (x_1 - \mu_1)(x_n - \mu_n) \\ (x_2 - \mu_2)(x_1 - \mu_1) & (x_2 - \mu_2)^2 & (x_2 - \mu_2)(x_3 - \mu_3) & \dots & (x_2 - \mu_2)(x_n - \mu_n) \\ (x_3 - \mu_3)(x_1 - \mu_1) & (x_3 - \mu_3)(x_2 - \mu_2) & (x_3 - \mu_3)^2 & \dots & (x_3 - \mu_3)(x_n - \mu_n) \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ (x_n - \mu_n)(x_1 - \mu_1) & (x_n - \mu_n)(x_2 - \mu_2) & (x_n - \mu_n)(x_3 - \mu_3) & \dots & (x_n - \mu_n)^2 \end{bmatrix}$$

با تعمیم تعریف وریانس با استفاده از عملگر امید ریاضی، برای یک بردار متغیر n بعدی $\underline{X}$ به ماتریس

وریانس-کوریانس $\underline{C}_x$ به صورت زیر خواهیم رسید.

$$\underline{C}_x = \begin{bmatrix} \sigma_1^2 & \sigma_{12} & \sigma_{13} & \dots & \sigma_{1n} \\ \sigma_{21} & \sigma_2^2 & \sigma_{23} & \dots & \sigma_{2n} \\ \sigma_{31} & \sigma_{32} & \sigma_3^2 & \dots & \sigma_{3n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \sigma_{n1} & \sigma_{n2} & \sigma_{n3} & \dots & \sigma_n^2 \end{bmatrix}$$

هر چند در این جزوه از ماتریس وزن و ماتریس کوفاکتور استفاده نمی شود، لیکن بنا بر تعریف رابطه بین

ماتریس وزن $\underline{W}_x$ ، ماتریس وریانس-کوریانس $\underline{C}_x$ و ماتریس کوفاکتور $\underline{Q}_x$ به صورت زیر برقرار می باشد.

$$\underline{W}_x = \sigma_0^2 \underline{C}_x^{-1}$$

$$\underline{Q}_x = \frac{1}{\sigma_0^2} \underline{C}_x$$

$$\underline{W}_x = \underline{Q}_x^{-1}$$

فصل سوم

انتشار خطاها

در مهندسی نقشه برداری نیز همانند بسیاری از رشته های علوم و مهندسی، کمیت هایی مانند طول و زاویه که مستقیماً مورد سنجش و اندازه گیری قرار می گیرند، غالباً برای محاسبه و برآورد کمیت های دیگری مانند مختصات نقاط می باشند. در چنین حالاتی کمیت های محاسباتی به صورت توابع ریاضی از اندازه گیری ها بیان می شوند. چنانچه اندازه گیری ها حاوی انواع خطاها باشند، کمیت های محاسبه شده از آنها نیز طبقاً همراه با مقادیری خطا خواهند بود. ارزیابی خطاها و بررسی رفتار و چگونگی انتشار آنها در کمیت های محاسبه شده به عنوان توابعی از خطاهای اندازه گیری، "انتشار خطاها" نامیده می شود که در این فصل به این مهم خواهیم پرداخت.

انتشار تابع توزیع

چنانچه در حالت یک بعدی تابعی پیوسته و مشتق پذیر مانند $y = g(x)$ داشته باشیم که x متغیر تصادفی با تابع چگالی $f_x(x)$ باشد، در آن صورت تابع چگالی $f_y(y)$ مطلوب می باشد. برای این منظور ابتدا تابع معکوس $y = g(x)$ را به صورت $x = h(y)$ بدست آورده و از رابطه زیر به تابع چگالی $f(y)$ می رسیم.

$$f_y(y) = f_x(h(y)) \left| \frac{\partial h(y)}{\partial y} \right| \quad (۱)$$

برای مثال تابع زیر را در نظر بگیرید.

$$y = g(x) = x^2, \quad f(x) = e^{-x^2/2}$$

حال اگر تابع معکوس را به صورت $\sqrt{y}$ در نظر بگیریم، تابع خواهیم داشت چگالی $f(y)$ بر اساس رابطه فوق بدست خواهد آمد.

$$f_y(y) = f_x(\sqrt{y}) \left| \frac{1}{2\sqrt{y}} \right| = e^{-y/2} \cdot \frac{1}{2\sqrt{y}}$$

در حالت چند بعدی نیز با تعمیم رابطه فوق برای بردار متغیرهای تصادفی $\underline{x}$ ، تابع چگالی $f_x(\underline{x})$ ، و بردار توابع $\underline{y} = \underline{g}(\underline{x})$ ، روابط زیر را خواهیم داشت.

$$\underline{x} = (x_1, x_2, x_3, \dots, x_n), \quad f_x(\underline{x}) = f_x(x_1, x_2, x_3, \dots, x_n)$$

$$\underline{y} = \underline{g}(\underline{x}) = g(x_1, x_2, x_3, \dots, x_n)$$

پس از تعیین توابع معکوس $\underline{h}(\underline{y})$ می توان تابع چگالی $f_y(\underline{y})$ را مانند زیر بدست آورد.

$$\underline{x} = \underline{h}(\underline{y}_1, \underline{y}_2, \underline{y}_3, \dots, \underline{y}_n)$$

$$f_y(\underline{y}) = f_x(h_1, h_2, h_3, \dots, h_n) |J_{hy}| \quad (2)$$

$$J_{hy} = \frac{\partial \underline{h}}{\partial \underline{y}} = \begin{bmatrix} \frac{\partial h_1}{\partial y_1} & \frac{\partial h_1}{\partial y_2} & \frac{\partial h_1}{\partial y_3} & \dots & \frac{\partial h_1}{\partial y_n} \\ \frac{\partial h_2}{\partial y_1} & \frac{\partial h_2}{\partial y_2} & \frac{\partial h_2}{\partial y_3} & \dots & \frac{\partial h_2}{\partial y_n} \\ \frac{\partial h_3}{\partial y_1} & \frac{\partial h_3}{\partial y_2} & \frac{\partial h_3}{\partial y_3} & \dots & \frac{\partial h_3}{\partial y_n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \frac{\partial h_n}{\partial y_1} & \frac{\partial h_n}{\partial y_2} & \frac{\partial h_n}{\partial y_3} & \dots & \frac{\partial h_n}{\partial y_n} \end{bmatrix}$$

برای مثال می توان دو تابع زیر را از متغیرهای تصادفی مستقل x_1 و x_2 با تابع چگالی نرمال استاندارد در نظر گرفت.

$$\begin{cases} y_1 = x_1 + x_2 \\ y_2 = \frac{x_1}{x_1 + x_2} \end{cases}$$

توابع معکوس روابط فوق به صورت زیر خواهد بود.

$$\begin{cases} h_1 = x_1 = y_1 y_2 \\ h_2 = x_2 = y_1 - y_1 y_2 \end{cases}$$

حال قدر مطلق ژاکوبین توابع معکوس را بدست می آوریم.

$$|J_{hy}| = \begin{vmatrix} y_2 & y_1 \\ 1-y_2 & -y_1 \end{vmatrix} = y_1$$

بنابراین تابع چگالی توام y_1 و y_2 به صورت زیر بدست خواهد آمد.

$$f_y(y_1, y_2) = f_x(h_1, h_2) |J_{hy}| = \frac{1}{\sqrt{2\pi}} e^{-\frac{(y_1 y_2)^2}{2}} \cdot \frac{1}{\sqrt{2\pi}} e^{-\frac{(y_1 - y_1 y_2)^2}{2}} \cdot y_1 = \frac{y_1}{2\pi} e^{-\frac{y_1^2 (2y_2 - 1)}{2}}$$

انتشار میانگین

میانگین واقعی هر کمیتی به کمک مفهوم امید ریاضی تعریف می شود. بر همین اساس با داشتن بردار متغیرهای تصادفی $\underline{x}$ ، بردار میانگین $\underline{\mu_x}$ و تابع چگالی $f_x(\underline{x})$ ، میانگین بردار توابع $\underline{y} = \underline{g}(\underline{x})$ به صورت زیر بدست خواهد آمد.

$$\begin{aligned} \underline{x} &= (x_1, x_2, x_3, \dots, x_n), \quad f_x(\underline{x}) = f_x(x_1, x_2, x_3, \dots, x_n) \\ \underline{\mu_x} &= (\mu_1, \mu_2, \mu_3, \dots, \mu_n) \\ \underline{y} = \underline{g}(\underline{x}) &= \underline{g}(x_1, x_2, x_3, \dots, x_n) \end{aligned}$$

$$E[y] = E[g(x_1, x_2, x_3, \dots, x_n)] = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} \dots \int_{-\infty}^{+\infty} g(x_1, x_2, x_3, \dots, x_n) f_x(x_1, x_2, x_3, \dots, x_n) dx_1 dx_2 dx_3 \dots dx_n$$

چنانچه انتشار میانگین را برای حالت خطی در نظر بگیریم، در آن صورت خواهیم داشت:

$$\underline{y} = \underline{A} \cdot \underline{x}$$

$$\underline{y} = (y_1, y_2, y_3, \dots, y_m)^T$$

$$\underline{x} = (x_1, x_2, x_3, \dots, x_n)^T$$

$$\underline{\mu}_x = (\mu_1, \mu_2, \mu_3, \dots, \mu_n)^T$$

$$\underline{A} = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1n} \\ a_{21} & a_{22} & a_{23} & \dots & a_{2n} \\ a_{31} & a_{32} & a_{33} & \dots & a_{3n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & a_{m3} & \dots & a_{mn} \end{bmatrix}$$

$$\underline{\mu}_y = E[\underline{y}] = E[\underline{A} \cdot \underline{x}] = \underline{A} \cdot E[\underline{x}] = \underline{A} \cdot \underline{\mu}_x \quad (3)$$

انتشار خطاهای سیستماتیک

چنانچه یک متغیر تصادفی x در نظر بگیریم و رابطه $y = f(x) = ax + b$ را داشته باشیم، در این صورت با فرض یک مقدار ثابت خطای سیستماتیک مانند Δx می خواهیم اثر آن را بر روی y که با نماد Δy معرفی می شود بدانیم.

$$y + \Delta y = f(x + \Delta x) = a(x + \Delta x) + b = ax + b + a\Delta x \quad (4)$$

با توجه به رابطه فوق در می یابیم که انتشار خطای سیستماتیک Δx در مدل داده شده برابر $a\Delta x$ است. یعنی

$$\Delta y = a\Delta x \quad (5)$$

حال اگر بخواهیم انتشار خطای سیستماتیک را برای بردار توابع خطی $\underline{y} = \underline{A}.\underline{x}$ ($\Delta \underline{y}$) با داشتن بردار متغیرهای تصادفی $\underline{x}$ ، بردار خطاهای سیستماتیک $\Delta \underline{x}$ بدست آوریم، مشابه حالت یک بعدی عمل می کنیم.

$$\underline{y} + \Delta \underline{y} = \underline{A}(\underline{x} + \Delta \underline{x}) = \underline{A}.\underline{x} + \underline{A}.\Delta \underline{x} \quad (6)$$

بنابراین انتشار بردار خطای سیستماتیک $\Delta \underline{x}$ در یک مدل خطی به صورت $\underline{A}.\Delta \underline{x}$ خواهد بود. در واقع از رابطه انتشار فوق هم برای تعیین مقدار اثر خطاهای سیستماتیک روی توابع مورد نظر، و هم برای اعمال تصحیحات مربوط به خطاهای سیستماتیک که برای هر اندازه گیری معلوم هستند می توان بهره برد. در صورتیکه انتشار خطاهای سیستماتیک را در یک مدل غیر خطی مورد ارزیابی قرار دهیم، با فرض کوچک بودن مقادیر خطاهای سیستماتیک به روابط زیر خواهیم رسید.

$$\begin{aligned} \underline{y} &= \underline{f}(\underline{x}) = \underline{f}(x_1, x_2, x_3, \dots, x_n) \\ \underline{y} + \Delta \underline{y} &= \underline{f}(\underline{x} + \Delta \underline{x}) = \underline{f}(x_1 + \Delta x_1, x_2 + \Delta x_2, x_3 + \Delta x_3, \dots, x_n + \Delta x_n) \end{aligned} \quad (7)$$

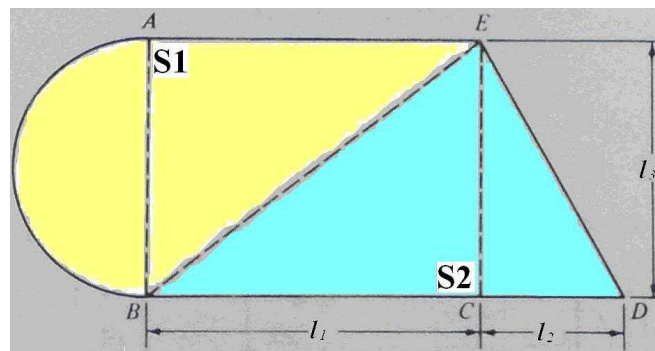
با بسط تیلور به خطی سازی مدل غیرخطی داده شده می پردازیم.

$$\underline{y} + \Delta \underline{y} = \underline{f}(x_1, x_2, x_3, \dots, x_n) + \frac{\partial f}{\partial x_1} \Delta x_1 + \frac{\partial f}{\partial x_2} \Delta x_2 + \frac{\partial f}{\partial x_3} \Delta x_3 + \dots + \frac{\partial f}{\partial x_n} \Delta x_n \quad (۸)$$

$$\Delta \underline{y} = \frac{\partial f}{\partial x_1} \Delta x_1 + \frac{\partial f}{\partial x_2} \Delta x_2 + \frac{\partial f}{\partial x_3} \Delta x_3 + \dots + \frac{\partial f}{\partial x_n} \Delta x_n = \frac{\partial f}{\partial \underline{x}} \Delta \underline{x}$$

$$= J_{yx} \Delta \underline{x} = \underline{A} \Delta \underline{x}$$

مثال زیر مربوط به مساحت دو قطعه زمین چسبیده به هم می باشد که در آن با اندازه گیری سه طول $l_1 = 50m$ ، $l_2 = 20m$ و $l_3 = 30m$ به ترتیب با خطاهای سیستماتیک $\delta l_1 = 2cm$ ، $\delta l_2 = -4cm$ و $\delta l_3 = 3cm$ مساحت های s_1 و s_2 قابل تعیین هستند. اثر خطاهای سیستماتیک طولیابی را بر روی دو مساحت s_1 و s_2 بدست آورید.



نگاره ۱- سه اندازه گیری طولی l_1 ، l_2 و l_3 برای تعیین مساحت دو قطعه $S1$ و $S2$

$$\begin{cases} s_1 = \frac{\pi}{8} l_3^2 + \frac{1}{2} l_1 l_3 \\ s_2 = \frac{1}{2} l_1 l_3 + \frac{1}{2} l_2 l_3 \end{cases}$$

ماتریس مشتقات جزئی روابط فوق به صورت زیر نوشته می شوند.

$$\begin{aligned}
 J_{sl} &= \begin{bmatrix} \frac{\partial s_1}{\partial l_1} & \frac{\partial s_1}{\partial l_2} & \frac{\partial s_1}{\partial l_3} \\ \frac{\partial s_2}{\partial l_1} & \frac{\partial s_2}{\partial l_2} & \frac{\partial s_2}{\partial l_3} \end{bmatrix} \\
 &= \begin{bmatrix} \frac{1}{2}l_3 & 0 & \left(\frac{1}{2}l_1 + \frac{\pi}{4}l_3\right) \\ \frac{1}{2}l_3 & \frac{1}{2}l_3 & \left(\frac{1}{2}l_1 + \frac{1}{2}l_3\right) \end{bmatrix} \\
 &= \begin{bmatrix} 15 & 0 & 49 \\ 15 & 15 & 35 \end{bmatrix}
 \end{aligned}$$

حال با استفاده از رابطه انتشار خطاهای سیستماتیک می توان اثر خطاهای سیستماتیک طولیابی را بر روی دو مساحت s_1 و s_2 بدست آورد.

$$\begin{aligned}
 \Delta \underline{s} &= J_{sl} \cdot \Delta \underline{l} \\
 \begin{bmatrix} \Delta s_1 \\ \Delta s_2 \end{bmatrix} &= \begin{bmatrix} 15 & 0 & 49 \\ 15 & 15 & 35 \end{bmatrix} \begin{bmatrix} 0.02 \\ -0.04 \\ 0.03 \end{bmatrix} \\
 \begin{bmatrix} \Delta s_1 \\ \Delta s_2 \end{bmatrix} &= \begin{bmatrix} 1.77^{m^2} \\ 0.75^{m^2} \end{bmatrix}
 \end{aligned}$$

انتشار خطاهای اتفاقی (انتشار وریانس کوریانس)

همانطور که قبلاً نیز مطرح شده است در مواجهه با مشاهدات آلوده به خطاهای بزرگ و سیستماتیک باید پس از شناسایی از مجموعه اندازه گیری ها حذف یا مورد اصلاح قرار گیرند. اما در برخورد با سایر خطاهای باقیمانده که دارای ماهیتی تصادفی می باشند و همواره در تمام اندازه گیری ها حضور دارند، هیچ چاره ای به جز اتخاذ روش های عملی و محاسباتی مناسب تر برای کاهش سطح آنها وجود ندارد. از طرفی اغلب اوقات این اندازه گیری ها به عنوان ورودی یکسری مدل های ریاضی منجر به کمیت های جدیدتری می شوند که دارای یک عدم قطعیت ناشی از خطاهای تصادفی اندازه گیری ها و مدل ریاضی مورد استفاده، می باشد. بنابراین بررسی رفتار خطاهای اتفاقی و چگونگی انتشار آنها در یک مدل ریاضی از جایگاه ویژه ای در تئوری خطاها برخوردار است. از آنجا که وریانس و کوریانس اندازه گیری ها بیانگر خطاهای اتفاقی موجود در اندازه

گیری است، انتشار خطاهای اتفاقی را می توان انتشار وریانس کوریانس نیز نامید. به همین دلیل این موضوع به قانون انتشار وریانس کوریانس معروف شده است.

برای سادگی درک انتشار وریانس کوریانس دو متغیر تصادفی X و Y را در نظر بگیرید که تواما n بار در حالیکه n به سمت $+\infty$ میل می کند، اندازه گیری شده اند.

$$\begin{pmatrix} X_i \\ Y_i \end{pmatrix} = \left[\begin{pmatrix} X_1 \\ Y_1 \end{pmatrix}, \begin{pmatrix} X_2 \\ Y_2 \end{pmatrix}, \begin{pmatrix} X_3 \\ Y_3 \end{pmatrix}, \dots, \begin{pmatrix} X_n \\ Y_n \end{pmatrix} \right] \quad (9)$$

اگر μ_x و μ_y به ترتیب میانگین واقعی X و Y باشند، در این صورت با فرض عدم وجود اشتباهات و خطاهای سیستماتیک، اختلافات آنها با اندازه گیری ها به عنوان خطاهای اتفاقی شناخته می شود.

$$\begin{pmatrix} r_{xi} \\ r_{yi} \end{pmatrix} = \left[\begin{pmatrix} X_1 - \mu_x \\ Y_1 - \mu_y \end{pmatrix}, \begin{pmatrix} X_2 - \mu_x \\ Y_2 - \mu_y \end{pmatrix}, \begin{pmatrix} X_3 - \mu_x \\ Y_3 - \mu_y \end{pmatrix}, \dots, \begin{pmatrix} X_n - \mu_x \\ Y_n - \mu_y \end{pmatrix} \right] \quad (10)$$

با استفاده از روابط زیر، می توان وریانس دو متغیر تصادفی X و Y و کوریانس بین آنها را بدست آورد.

$$\sigma_x^2 = \frac{1}{n} \sum_{i=1}^n r_{xi}^2 \quad (11)$$

$$\sigma_y^2 = \frac{1}{n} \sum_{i=1}^n r_{yi}^2 \quad (12)$$

$$\sigma_{xy} = \frac{1}{n} \sum_{i=1}^n r_{xi} r_{yi} \quad (13)$$

حال چنانچه دو تابع F و G را بر حسب متغیرهای تصادفی X و Y به صورت زیر تعریف نماییم، می خواهیم نحوه اثر گذاری وریانس و کوریانس دو متغیر تصادفی X و Y را در دو تابع F و G بررسی نماییم. در واقع می خواهیم وریانس و کوریانس دو تابع F و G را بدست آوریم.

$$\begin{aligned} F &= a_{11}X + a_{12}Y + b_1 \\ G &= a_{21}X + a_{22}Y + b_2 \end{aligned} \quad (14)$$

با بدست آوردن میانگین F و G ، یعنی μ_f و μ_g ، می توان با توجه به رابطه فوق اختلافات F و G را نسبت آنها تعیین نمود.

$$\begin{aligned} r_{fi} &= a_{11}r_{xi} + a_{12}r_{yi} \\ r_{gi} &= a_{21}r_{xi} + a_{22}r_{yi} \end{aligned} \quad (15)$$

حال مشابه متغیرهای تصادفی X و Y ، وریانس و کوریانس دو تابع F و G نیز به صورت زیر قابل تعیین است.

$$\sigma_f^2 = \frac{1}{n} \sum_{i=1}^n r_{fi}^2 \quad (16)$$

$$\sigma_g^2 = \frac{1}{n} \sum_{i=1}^n r_{gi}^2 \quad (17)$$

$$\sigma_{fg} = \frac{1}{n} \sum_{i=1}^n r_{fi} r_{gi} \quad (18)$$

چنانچه در سمت راست رابطه وریانس تابع F ، معادل انحرافات یا باقیمانده های r_{fi} را بر حسب r_{xi} و r_{yi} بنویسیم، به رابطه زیر خواهیم رسید.

$$\begin{aligned}\sigma_f^2 &= \frac{1}{n} \sum_{i=1}^n (a_{11}r_{xi} + a_{12}r_{yi})^2 \\ &= \frac{1}{n} \sum_{i=1}^n (a_{11}^2 r_{xi}^2 + 2a_{11}a_{12}r_{xi}r_{yi} + a_{12}^2 r_{yi}^2) \\ &= \frac{1}{n} \sum_{i=1}^n a_{11}^2 r_{xi}^2 + \frac{1}{n} \sum_{i=1}^n 2a_{11}a_{12}r_{xi}r_{yi} + \frac{1}{n} \sum_{i=1}^n a_{12}^2 r_{yi}^2 \\ &= a_{11}^2 \sigma_x^2 + 2a_{11}a_{12} \sigma_{xy} + a_{12}^2 \sigma_y^2\end{aligned}\quad (19)$$

به طور مشابه برای وریانس تابع G و کوریانس بین F و G به روابط زیر خواهیم رسید.

$$\sigma_g^2 = a_{21}^2 \sigma_x^2 + 2a_{21}a_{22} \sigma_{xy} + a_{22}^2 \sigma_y^2 \quad (20)$$

$$\sigma_{fg} = a_{11}a_{21} \sigma_x^2 + (a_{11}a_{22} + a_{12}a_{21}) \sigma_{xy} + a_{12}a_{22} \sigma_y^2 \quad (21)$$

حال فرم ماتریسی مربوط به دو تابع F و G و وریانس و کوریانس متغیرهای تصادفی X و Y را در نظر می گیریم.

$$\underline{H} = \begin{bmatrix} F \\ G \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \begin{bmatrix} X \\ Y \end{bmatrix} + \begin{bmatrix} b_1 \\ b_2 \end{bmatrix} = \underline{AZ} + \underline{b} \quad (22)$$

$$C_z = \begin{bmatrix} \sigma_x^2 & \sigma_{xy} \\ \sigma_{xy} & \sigma_y^2 \end{bmatrix} \quad (23)$$

با توجه به جواب های بدست آمده برای وریانس و کوریانس دو تابع F و G متوجه می شویم که رابطه زیر برقرار است.

$$C_h = \begin{bmatrix} \sigma_f^2 & \sigma_{fg} \\ \sigma_{fg} & \sigma_g^2 \end{bmatrix} = \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix} \begin{bmatrix} \sigma_x^2 & \sigma_{xy} \\ \sigma_{xy} & \sigma_y^2 \end{bmatrix} \begin{bmatrix} a_{11} & a_{12} \\ a_{21} & a_{22} \end{bmatrix}^T = \underline{A} \underline{C}_z \underline{A}^T \quad (24)$$

اکنون با توجه به مثال فوق به حل کلی انتشار وریانس کوریانس برای دو مدل خطی و غیرخطی می پردازیم.

انتشار وریانس کوریانس در مدل های خطی

یک بردار متغیر تصادفی n تایی $\underline{X}$ را همراه بردار میانگین $\underline{\mu}_x$ و ماتریس وریانس کوریانس $\underline{C}_x$ آن در اختیار داریم و m تابع خطی Y_i را مانند زیر در نظر می گیریم.

$$\underline{X} = (X_1, X_2, X_3, \dots, X_n); \quad \underline{\mu}_x = (\mu_1, \mu_2, \mu_3, \dots, \mu_n); \quad \underline{C}_x \quad (25)$$

$$\begin{cases} Y_1 = a_{11}X_1 + a_{12}X_2 + a_{13}X_3 + \dots + a_{1n}X_n \\ Y_2 = a_{21}X_1 + a_{22}X_2 + a_{23}X_3 + \dots + a_{2n}X_n \\ Y_3 = a_{31}X_1 + a_{32}X_2 + a_{33}X_3 + \dots + a_{3n}X_n \\ \vdots \quad \quad \quad \vdots \quad \quad \quad \dots \quad \quad \vdots \\ Y_m = a_{m1}X_1 + a_{m2}X_2 + a_{m3}X_3 + \dots + a_{mn}X_n \end{cases}$$

نمایش ماتریسی توابع خطی Y_i به صورت زیر خواهد بود.

$$\underline{Y} = \underline{A} \underline{X} \quad (26)$$

که بسط اجزای آن به شرح زیر می باشد.

$$\underline{Y} = \begin{bmatrix} Y_1 \\ Y_2 \\ Y_3 \\ \vdots \\ Y_m \end{bmatrix}; \quad \underline{A} = \begin{bmatrix} a_{11} & a_{12} & a_{13} & \dots & a_{1n} \\ a_{21} & a_{22} & a_{23} & \dots & a_{2n} \\ a_{31} & a_{32} & a_{33} & \dots & a_{3n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ a_{m1} & a_{m2} & a_{m3} & \dots & a_{mn} \end{bmatrix}; \quad \underline{X} = \begin{bmatrix} X_1 \\ X_2 \\ X_3 \\ \vdots \\ X_n \end{bmatrix}$$

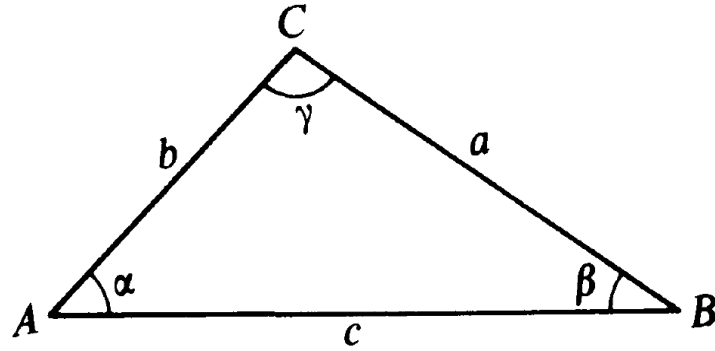
بردار میانگین توابع خطی $\underline{Y}$ بر پایه امید ریاضی به صورت زیر بدست می آید.

$$\underline{\mu}_y = E[\underline{Y}] = E[\underline{AX}] = \underline{AE}[\underline{X}] = \underline{A}\underline{\mu}_x \quad (27)$$

حال با استفاده از تعریف وریانس کوریانس بر پایه امید ریاضی روابط زیر را تنظیم و سپس به ماتریس وریانس کوریانس توابع خطی $\underline{Y}$ می رسمیم.

$$\begin{aligned} \underline{C}_y &= E[(\underline{Y} - \underline{\mu}_y)(\underline{Y} - \underline{\mu}_y)^T] \\ &= E[(\underline{AX} - \underline{A}\underline{\mu}_x)(\underline{AX} - \underline{A}\underline{\mu}_x)^T] \\ &= E[\underline{A}(\underline{X} - \underline{\mu}_x)(\underline{X} - \underline{\mu}_x)^T \underline{A}^T] \\ &= \underline{AE}[(\underline{X} - \underline{\mu}_x)(\underline{X} - \underline{\mu}_x)^T] \underline{A}^T \\ &= \underline{AC}_x \underline{A}^T \end{aligned} \quad (28)$$

برای مثال مثلث ABC را در نگاره (۲) در نظر بگیرید که در آن سه زاویه α ، β و γ با دقت های برابر اندازه گیری شده اند. چنانچه اندازه گیری ها مستقل از هم باشند و انحراف معیار اندازه گیری هر امتداد s_d باشد، خطای مجاز بست مثلث ABC را با استفاده از روابط فوق بدست آورید.



نگاره ۲- مثلث ABC با سه زاویه اندازه گیری شده α ، β و γ

ابتدا دقت هر زاویه را بدست می آوریم.

$$\alpha = d_2 - d_1$$

$$s_\alpha^2 = s_\beta^2 = s_\gamma^2 = \begin{bmatrix} -1 & 1 \end{bmatrix} \begin{bmatrix} s_d^2 & \\ & s_d^2 \end{bmatrix} \begin{bmatrix} -1 \\ 1 \end{bmatrix} = 2s_d^2$$

حال دقت خطای بست زاویه مثلث که همان خطای مجاز زاویه ای مثلث است را بدست می آوریم.

$$\delta = \alpha + \beta + \gamma - \pi$$

$$s_\delta^2 = \begin{bmatrix} 1 & 1 & 1 \end{bmatrix} \begin{bmatrix} s_\alpha^2 & & \\ & s_\alpha^2 & \\ & & s_\alpha^2 \end{bmatrix} \begin{bmatrix} 1 \\ 1 \\ 1 \end{bmatrix} = 3s_\alpha^2 = 6s_d^2$$

انتشار وریانس کوریانس در مدل های غیر خطی

یک بردار متغیر تصادفی n تایی $\underline{X}$ را همراه بردار میانگین $\underline{\mu}_x$ و ماتریس وریانس کوریانس $\underline{C}_x$ آن در

اختیار داریم و m تابع غیر خطی Y_i را مانند زیر در نظر می گیریم.

$$\underline{X} = (X_1, X_2, X_3, \dots, X_n); \quad \underline{\mu}_x = (\mu_1, \mu_2, \mu_3, \dots, \mu_n); \quad \underline{C}_x \quad (29)$$

$$\begin{cases} Y_1 = f_1(X_1, X_2, X_3, \dots, X_n) \\ Y_2 = f_2(X_1, X_2, X_3, \dots, X_n) \\ Y_3 = f_3(X_1, X_2, X_3, \dots, X_n) \\ \vdots \\ Y_m = f_m(X_1, X_2, X_3, \dots, X_n) \end{cases}$$

نمایش ماتریسی توابع غیر خطی Y_i به صورت زیر خواهد بود.

$$\underline{Y} = \underline{f}(\underline{X}) \quad (30)$$

که بسط اجزای آن به شرح زیر می باشد.

$$\underline{Y} = \begin{bmatrix} Y_1 \\ Y_2 \\ Y_3 \\ \vdots \\ Y_m \end{bmatrix}; \quad \underline{f} = \begin{bmatrix} f_1 \\ f_2 \\ f_3 \\ \vdots \\ f_m \end{bmatrix}; \quad \underline{X} = \begin{bmatrix} X_1 \\ X_2 \\ X_3 \\ \vdots \\ X_n \end{bmatrix}$$

بردار میانگین توابع خطی $\underline{Y}$ بر پایه امید ریاضی به صورت زیر بدست می آید.

$$\underline{\mu}_y = E[\underline{Y}] = E[\underline{f}(\underline{X})] \quad (31)$$

حال با استفاده از بسط تیلور حول بردار میانگین $\underline{\mu}_x$ ، توابع غیر خطی $\underline{Y}$ را به صورت زیر خطی می کنیم.

$$\underline{Y} = \underline{f}(\underline{\mu}_x) + \frac{\partial \underline{f}}{\partial \underline{X}} (\underline{X} - \underline{\mu}_x) + \dots \quad (32)$$

با صرف نظر کردن از ترم های بالاتر از درجه اول و بسط رابطه فوق به رابطه زیر می رسیم.

$$\left\{ \begin{array}{l} Y_1 = f_1(\underline{\mu}_x) + \frac{\partial f_1}{\partial x_1}(X_1 - \mu_{x1}) + \frac{\partial f_1}{\partial x_2}(X_2 - \mu_{x2}) + \frac{\partial f_1}{\partial x_3}(X_3 - \mu_{x3}) + \dots + \frac{\partial f_1}{\partial x_n}(X_n - \mu_{xn}) \\ Y_2 = f_2(\underline{\mu}_x) + \frac{\partial f_2}{\partial x_1}(X_1 - \mu_{x1}) + \frac{\partial f_2}{\partial x_2}(X_2 - \mu_{x2}) + \frac{\partial f_2}{\partial x_3}(X_3 - \mu_{x3}) + \dots + \frac{\partial f_2}{\partial x_n}(X_n - \mu_{xn}) \\ Y_3 = f_3(\underline{\mu}_x) + \frac{\partial f_3}{\partial x_1}(X_1 - \mu_{x1}) + \frac{\partial f_3}{\partial x_2}(X_2 - \mu_{x2}) + \frac{\partial f_3}{\partial x_3}(X_3 - \mu_{x3}) + \dots + \frac{\partial f_3}{\partial x_n}(X_n - \mu_{xn}) \\ \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \quad \quad \quad \vdots \\ Y_m = f_m(\underline{\mu}_x) + \frac{\partial f_m}{\partial x_1}(X_1 - \mu_{x1}) + \frac{\partial f_m}{\partial x_2}(X_2 - \mu_{x2}) + \frac{\partial f_m}{\partial x_3}(X_3 - \mu_{x3}) + \dots + \frac{\partial f_m}{\partial x_n}(X_n - \mu_{xn}) \end{array} \right. \quad (33)$$

در فرم ماتریسی بسط تیلور توابع غیر خطی $\underline{Y}$ ، از بردار ها و ماتریس مشتقات جزئی (ماتریس ژاکوبین) به شرح زیر استفاده می کنیم.

$$\underline{Y} = \begin{bmatrix} Y_1 \\ Y_2 \\ Y_3 \\ \vdots \\ Y_m \end{bmatrix}; \quad \underline{f}(\underline{\mu}_x) = \begin{bmatrix} f_1(\underline{\mu}_x) \\ f_2(\underline{\mu}_x) \\ f_3(\underline{\mu}_x) \\ \vdots \\ f_m(\underline{\mu}_x) \end{bmatrix}; \quad \underline{J}_{yx} = \frac{\partial \underline{f}}{\partial \underline{X}} = \begin{bmatrix} \frac{\partial f_1}{\partial x_1} & \frac{\partial f_1}{\partial x_2} & \frac{\partial f_1}{\partial x_3} & \dots & \frac{\partial f_1}{\partial x_n} \\ \frac{\partial f_2}{\partial x_1} & \frac{\partial f_2}{\partial x_2} & \frac{\partial f_2}{\partial x_3} & \dots & \frac{\partial f_2}{\partial x_n} \\ \frac{\partial f_3}{\partial x_1} & \frac{\partial f_3}{\partial x_2} & \frac{\partial f_3}{\partial x_3} & \dots & \frac{\partial f_3}{\partial x_n} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \frac{\partial f_n}{\partial x_1} & \frac{\partial f_n}{\partial x_2} & \frac{\partial f_n}{\partial x_3} & \dots & \frac{\partial f_n}{\partial x_n} \end{bmatrix}; \quad (\underline{X} - \underline{\mu}_x) = \begin{bmatrix} X_1 - \mu_{x1} \\ X_2 - \mu_{x2} \\ X_3 - \mu_{x3} \\ \vdots \\ X_n - \mu_{xn} \end{bmatrix}$$

حال از طرفین رابطه (۳۲) امید ریاضی می گیریم و نتیجه آن را از همان رابطه کم می کنیم.

$$\underline{Y} - \underline{\mu}_y = \frac{\partial \underline{f}}{\partial \underline{X}}(\underline{X} - \underline{\mu}_x) = \underline{J}_{yx}(\underline{X} - \underline{\mu}_x) \quad (34)$$

مطابق تعریف وریانس کوریانس بر پایه امید ریاضی و توجه به رابطه فوق طی مراحل زیر به قانون انتشار

وریانس کوریانس در مدل های غیر خطی دست می یابیم.

$$\begin{aligned}
 \underline{C}_y &= E[(\underline{Y} - \underline{\mu}_y)(\underline{Y} - \underline{\mu}_y)^T] \\
 &= E[\underline{J}_{yx}(\underline{X} - \underline{\mu}_x)(\underline{X} - \underline{\mu}_x)^T \underline{J}_{yx}^T] \\
 &= \underline{J}_{yx} E[(\underline{X} - \underline{\mu}_x)(\underline{X} - \underline{\mu}_x)^T] \underline{J}_{yx}^T \\
 &= \underline{J}_{yx} \underline{C}_x \underline{J}_{yx}^T
 \end{aligned} \tag{۳۵}$$

برای مثال مجدداً مثلث ABC را در نگاره (۲) در نظر بگیرید. این بار زوایای α و β و طول c مورد اندازه گیری قرار گرفته اند. با فرض استقلال بین اندازه گیری ها و در نظر گرفتن انحراف معیارهای σ_α ، σ_β و σ_c به ترتیب برای زوایای α و β و طول c ، انحراف معیار طول a مطلوب می باشد.

$$a = \frac{c \sin \alpha}{\sin \gamma} = \frac{c \sin \alpha}{\sin(\pi - \alpha - \beta)}$$

$$Y = a = f(\underline{l})$$

$$\underline{l} = (\alpha, \beta, c)^T$$

$$\underline{C}_l = \begin{bmatrix} \sigma_\alpha^2 & 0 & 0 \\ 0 & \sigma_\beta^2 & 0 \\ 0 & 0 & \sigma_c^2 \end{bmatrix}$$

$$\begin{aligned}
 \underline{J}_{al} &= \begin{bmatrix} \frac{\partial a}{\partial \alpha} & \frac{\partial a}{\partial \beta} & \frac{\partial a}{\partial c} \end{bmatrix} \\
 &= \begin{bmatrix} \frac{c(\cos \alpha \sin \gamma \sin \alpha \cos \gamma)}{\sin^2 \gamma} & \frac{c(\sin \alpha \cos \gamma)}{\sin^2 \gamma} & \frac{\sin \alpha}{\sin \gamma} \end{bmatrix} \\
 &= \begin{bmatrix} (a \cot \alpha + a \cot \gamma) & a \cot \gamma & \frac{a}{c} \end{bmatrix}
 \end{aligned}$$

$$\sigma_a^2 = \underline{C}_a = \underline{J}_{al} \underline{C}_l \underline{J}_{al}^T$$

$$\sigma_a^2 = a^2 (\cot \alpha + \cot \gamma)^2 \sigma_\alpha^2 + a^2 \cot^2 \gamma \sigma_\beta^2 + \frac{a^2}{c^2} \sigma_c^2$$

انتشار ماتریس کوریانس (مبحث تکمیلی)

برای شروع مطلب سه بردار $\underline{y}$ ، $\underline{z}$ و $\underline{r}$ را به عنوان توابعی از بردار متغیر تصادفی $\underline{x}$ با ماتریس وریانس کوریانس $\underline{C}_x$ به صورت زیر در نظر بگیرید.

$$\begin{aligned}\underline{y} &= \underline{A}\underline{x} + \underline{a} \\ \underline{z} &= \underline{B}\underline{y} + \underline{b} \\ \underline{r} &= \underline{C}\underline{z} + \underline{c}\end{aligned}\tag{۳۶}$$

علاوه بر یافتن ماتریس های وریانس کوریانس $\underline{C}_y$ ، $\underline{C}_z$ و $\underline{C}_r$ می خواهیم ماتریس های وابستگی $\underline{C}_{yz}$ ، $\underline{C}_{yr}$ و $\underline{C}_{zr}$ را نیز بدست آوریم. ابتدا بر اساس قانون انتشار وریانس کوریانس مطابق زیر ماتریس های وریانس کوریانس $\underline{C}_y$ ، $\underline{C}_z$ و $\underline{C}_r$ بدست می آوریم.

$$\begin{aligned}\underline{C}_y &= \underline{J}_{yx} \underline{C}_x \underline{J}_{yx}^T = \underline{A} \underline{C}_x \underline{A}^T \\ \underline{C}_z &= \underline{J}_{zy} \underline{C}_y \underline{J}_{zy}^T = \underline{B} \underline{C}_y \underline{B}^T = \underline{B} \underline{A} \underline{C}_x \underline{A}^T \underline{B}^T \\ \underline{C}_r &= \underline{J}_{rz} \underline{C}_z \underline{J}_{rz}^T = \underline{C} \underline{C}_z \underline{C}^T = \underline{C} \underline{B} \underline{A} \underline{C}_x \underline{A}^T \underline{B}^T \underline{C}^T\end{aligned}\tag{۳۷}$$

حال برای تعیین ماتریس وابستگی $\underline{C}_{yz}$ با بهره گیری از قانون انتشار وریانس کوریانس، به رابطه زیر می رسیم.

$$\underline{C}_{yz} = \underline{J}_{yx} \underline{C}_{xy} \underline{J}_{zy}^T = \underline{A} \underline{C}_{xy} \underline{B}^T\tag{۳۸}$$

اما در این رابطه ماتریس کوریانس $\underline{C}_{xy}$ مجهول می باشد. برای رفع این مشکل روابط زیر را تنظیم می کنیم.

$$\begin{aligned}\underline{x} &= \underline{I} \underline{x} \\ \underline{y} &= \underline{A} \underline{x} + \underline{a}\end{aligned}\tag{۳۹}$$

حال مجدداً با توجه به قانون انتشار وریانس کوریانس، می توان ماتریس کوریانس $\underline{C}_{xy}$ را بدست آورد.

$$\underline{C}_{xy} = \underline{J}_{xx} \underline{C}_x \underline{J}_{yx}^T = \underline{I} \underline{C}_x \underline{A}^T = \underline{C}_x \underline{A}^T\tag{۴۰}$$

اکنون با جایگذاری رابطه فوق در (۳۸) به ماتریس وابستگی $\underline{C}_{yz}$ دست می یابیم.

$$\underline{C}_{yz} = \underline{A} \underline{C}_{xy} \underline{B}^T = \underline{A} \underline{C}_x \underline{A}^T \underline{B}^T\tag{۴۱}$$

برای تعیین ماتریس وابستگی $\underline{C}_{yr}$ مشابه $\underline{C}_{yz}$ به رابطه زیر می رسم.

$$\underline{C}_{yr} = \underline{J}_{yx} \underline{C}_{xz} \underline{J}_{rz}^T = \underline{A} \underline{C}_{xz} \underline{C}^T\tag{۴۲}$$

مجدداً با نگاه به رابطه فوق در می یابیم که $\underline{C}_{xz}$ مجهول است. مشابه قبل برای رفع مشکل مذکور به روابط زیر را تنظیم می کنیم.

$$\begin{aligned}\underline{x} &= \underline{I} \underline{x} \\ \underline{z} &= \underline{B} \underline{y} + \underline{b}\end{aligned}\tag{۴۳}$$

حال می توان ماتریس کوریانس $\underline{C}_{xz}$ را بر اساس قانون انتشار وریانس کوریانس و توجه به رابطه (۴۰) بدست آورد.

$$\underline{C}_{xz} = \underline{J}_{xx} \underline{C}_{xy} \underline{J}_{zy}^T = \underline{I} \underline{C}_{xy} \underline{B}^T = \underline{C}_x \underline{A}^T \underline{B}^T \quad (44)$$

اکنون به رابطه (۴۲) بر می گردیم و با استفاده از (۴۴) ماتریس وابستگی $\underline{C}_{yr}$ را به صورت زیر بدست می آوریم.

$$\underline{C}_{yr} = \underline{A} \underline{C}_{xz} \underline{C}^T = \underline{A} \underline{C}_x \underline{A}^T \underline{B}^T \underline{C}^T \quad (45)$$

مطابق شیوه فوق ماتریس کوریانس $\underline{C}_{zr}$ نیز به صورت زیر قابل تعیین است.

$$\underline{C}_{zr} = \underline{J}_{zy} \underline{C}_{yz} \underline{J}_{rz}^T = \underline{B} \underline{C}_{yz} \underline{C}^T \quad (46)$$

$\underline{C}_{yz}$ قبلا در رابطه (۴۱) بدست آمده است و با جایگذاری آن در رابطه فوق به جواب نهایی برای ماتریس کوریانس $\underline{C}_{zr}$ می رسیم.

$$\underline{C}_{zr} = \underline{B} \underline{C}_{yz} \underline{C}^T = \underline{B} \underline{A} \underline{C}_x \underline{A}^T \underline{B}^T \underline{C}^T \quad (47)$$

به این ترتیب طی مثال فوق علاوه بر بررسی مجدد انتشار وریانس کوریانس به موضوع جدید انتشار کوریانس نیز پرداخته شد.

حال برای بررسی جامع و کلی انتشار کوریانس دو بردار متغیر تصادفی مانند $\underline{x}$ و $\underline{t}$ را در نظر بگیرید که به یکدیگر وابسته هستند و ماتریس های وریانس کوریانس هریک و ماتریس وابستگی آنها را به ترتیب با $\underline{C}_x$ ،

$\underline{C}_{xt}$ و $\underline{C}_t$ نمایش داده می شوند. چنانچه $\underline{y}$ و $\underline{z}$ را به ترتیب به عنوان توابعی از $\underline{x}$ و $\underline{t}$ در نظر بگیریم، می خواهیم ماتریس وابستگی بین $\underline{y}$ و $\underline{z}$ یعنی $\underline{C}_{yz}$ را بدست آوریم.

$$\begin{aligned}\underline{y} &= \underline{f}(\underline{x}) \\ \underline{z} &= \underline{g}(\underline{t})\end{aligned}\quad (48)$$

ابتدا دو بردار $\underline{r}$ و $\underline{s}$ را به صورت زیر تعریف می کنیم.

$$\underline{r} = \begin{bmatrix} \underline{y} \\ \underline{z} \end{bmatrix}, \quad \underline{s} = \begin{bmatrix} \underline{x} \\ \underline{t} \end{bmatrix}\quad (49)$$

با توجه به اینکه $\underline{y}$ و $\underline{z}$ خود توابعی از $\underline{x}$ و $\underline{t}$ هستند، بنابراین می توان $\underline{r}$ را تابعی از $\underline{s}$ دانست.

$$\underline{r} = \underline{F}(\underline{s})\quad (50)$$

حال قانون انتشار وریانس کوریانس را برای بردار $\underline{r}$ که تابعی از بردار $\underline{s}$ است اعمال می کنیم.

$$\begin{aligned}\underline{C}_r &= \begin{bmatrix} \underline{C}_y & \underline{C}_{yz} \\ \underline{C}_{zy} & \underline{C}_z \end{bmatrix} = \underline{J}_{rs} \underline{C}_s \underline{J}_{rs}^T \\ &= \begin{bmatrix} \underline{J}_{yx} & \underline{J}_{yt} \\ \underline{J}_{zx} & \underline{J}_{zt} \end{bmatrix} \begin{bmatrix} \underline{C}_x & \underline{C}_{xt} \\ \underline{C}_{tx} & \underline{C}_t \end{bmatrix} \begin{bmatrix} \underline{J}_{yx} & \underline{J}_{yt} \\ \underline{J}_{zx} & \underline{J}_{zt} \end{bmatrix}^T \\ &= \begin{bmatrix} \underline{J}_{yx} & \underline{0} \\ \underline{0} & \underline{J}_{zt} \end{bmatrix} \begin{bmatrix} \underline{C}_x & \underline{C}_{xt} \\ \underline{C}_{tx} & \underline{C}_t \end{bmatrix} \begin{bmatrix} \underline{J}_{yx} & \underline{0} \\ \underline{0} & \underline{J}_{zt} \end{bmatrix}^T \\ &= \begin{bmatrix} \underline{J}_{yx} \underline{C}_x \underline{J}_{yx}^T & \underline{J}_{yx} \underline{C}_{xt} \underline{J}_{zt}^T \\ \underline{J}_{zt} \underline{C}_{tx} \underline{J}_{yx}^T & \underline{J}_{zt} \underline{C}_t \underline{J}_{zt}^T \end{bmatrix}\end{aligned}\quad (51)$$

با توجه به ساختار ماتریس وریانس کوریانس $\underline{C}_r$ ، می توان روابط زیر را استخراج نمود.

$$\underline{C}_y = \underline{J}_{yx} \underline{C}_x \underline{J}_{yx}^T \quad (52)$$

$$\underline{C}_{yz} = \underline{J}_{yx} \underline{C}_{xt} \underline{J}_{zt}^T$$

$$\underline{C}_{zy} = \underline{J}_{zt} \underline{C}_{tx} \underline{J}_{yx}^T = \underline{C}_{yz}^T$$

$$\underline{C}_z = \underline{J}_{zt} \underline{C}_t \underline{J}_{zt}^T$$

بنابراین از روابط فوق پیداست که چگونه می توان ماتریس کوریانس دو بردار $\underline{y}$ و $\underline{z}$ را برحسب ماتریس

کوریانس دو بردار $\underline{x}$ و $\underline{t}$ بدست آورد. $(\underline{C}_{xt})$

فصل چہارم

نمونه برداری و برآورد

از آنجا که معمولاً ما با مجموعه ای بسیار بزرگ یا نامحدود از تعداد حالات ممکن برای یک متغیر تصادفی مواجه هستیم، ناگزیر به انتخاب زیر مجموعه ای کوچکتر از آن می‌باشیم یا بتوانیم به صورت عملی به برآوردهایی از کمیت‌های مورد نظر دست یابیم. بنابراین در این فصل ضمن تشریح مبانی نمونه برداری و برآورد به برخی برآوردهای مهم از جمله برآورد میانگین وزندار پرداخته خواهد شد.

جمعیت و نمونه

جمعیت یا جامعه به مجموعه تمام حالات ممکن برای یک یا چند متغیر تصادفی اعم از اینکه محدود یا نامحدود باشد اطلاق می‌گردد. در واقع زمانی که با اندازه‌گیری سروکار داریم، به مجموعه کل اندازه‌گیری‌های ممکن که می‌تواند شامل بی‌نهایت تکرار باشد، جمعیت گفته می‌شود. در مقابل تعریف جمعیت، به زیر مجموعه ای از جمعیت که قسمت کوچکی از کل اندازه‌گیری‌ها می‌باشد، نمونه اطلاق می‌گردد. بنابر تعریف اخیر، هر مجموعه ای از اندازه‌گیری‌های نقشه برداری بیانگر یک نمونه خواهد بود.

برآورد

از آنجا که اغلب اوقات دستیابی به تمام عناصر یک جامعه غیر ممکن یا بسیار پرهزینه است، ما با کمک نمونه‌ها سعی در برآورد یا تخمین پارامترهای مجهول جامعه مورد مطالعه داریم. این برآورد‌ها عموماً به دو دسته نقطه‌ای و فاصله‌ای تقسیم بندی می‌شوند. چنانچه نتیجه برآوردها مانند میانگین و واریانس به یک مقدار معین اشاره نماید، برآورد را نقطه‌ای می‌نامند. چنانچه عدم قطعیت برآوردها را نیز در نظر بگیریم و در واقع برآورد خود را از یک پارامتر مجهول در یک بازه نمایش دهیم، در این صورت به آن برآورد فاصله‌ای می‌گوییم. آماره ای که برای برآورد نقطه‌ای مورد استفاده قرار می‌گیرد، برآوردگر یا تخمین‌گر نامیده می‌شود. در واقع هر برآوردگر یک مدل ریاضی است که چگونگی محاسبه یک برآورد را از یک نمونه مشخص می‌

سازد. به عنوان مثال برای یک نمونه n تایی از متغیر تصادفی x ، برآورد میانگین $(\bar{x})$ و وریانس (s_x^2) جامعه از روی نمونه موجود به صورت زیر می باشند.

$$\bar{x} = \frac{1}{n} \sum_{i=1}^n x_i \quad (1)$$

$$s_x^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 \quad (2)$$

معیار های کیفیت برآوردگرها

از آنجا که روش های برآورد متفاوتی برای هر کمیت وجود دارد، باید یک برآوردگر خوب تا حد امکان به مقدار واقعی پارامتر نزدیک باشد. از این رو به منظور ارزیابی کیفیت میزان اعتماد پذیری یک برآوردگر معیارهای زیر معرفی می شوند.

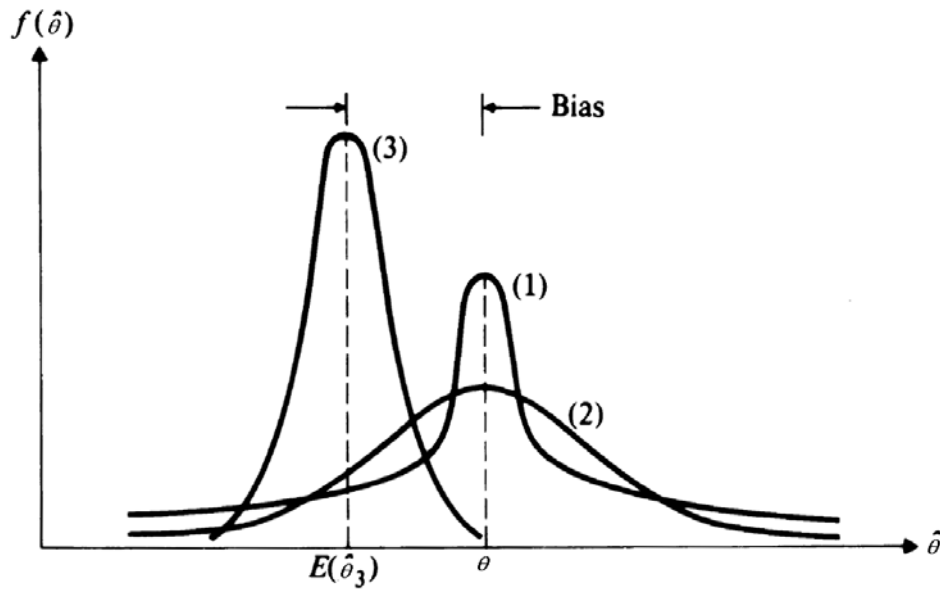
ناریبی

چنانچه برای برآورد یک پارامتر مجهول مانند θ از یک برآوردگر مانند $\hat{\theta}$ استفاده شود، در صورتی که حجم نمونه بسیار بزرگ باشد و امید ریاضی آن برابر با مقدار واقعی پارامتر باشد به آن برآوردگر ناریب اطلاق می شود.

$$E[\hat{\theta}] = \theta \quad (3)$$

ویژگی ناریبی یکی ویژگی ضروری برای برآوردگرهای خوب به حساب می آید. برآوردگر های میانگین و وریانس دارای این ویژگی می باشند و بنابراین برآوردگرهایی ناریب هستند. برای مثال در نگاره (۱) توزیع سه

برآوردگر متفاوت برای یک کمیت نشان داده شده است. همانطور که در این نگاره پیداست دو برآوردگر $\hat{\theta}_1$ و $\hat{\theta}_2$ ناریب ولی برآوردگر $\hat{\theta}_3$ هر چند دقیق تر است دارای اریب می باشد.



نگاره ۱- مقایسه برآوردگر اریب و ناریب

سازگاری

چنانچه برای برآورد یک پارامتر مجهول مانند θ از یک برآوردگر مانند $\hat{\theta}$ استفاده شود، در صورتی که حجم نمونه بسیار بزرگ شود و وریانس برآوردگر ناریب به سمت صفر میل کند به آن برآوردگر سازگار اطلاق می شود.

$$\lim_{n \rightarrow \infty} \text{var}[\hat{\theta}] = 0 \quad (۴)$$

این ویژگی برای میانگین هر نمونه بزرگ وجود دارد.

$$\text{var}(\bar{x}_n) = \frac{\sigma_x^2}{n} \quad (۵)$$

$$\lim_{n \rightarrow \infty} \text{var}(\bar{x}_n) = \lim_{n \rightarrow \infty} \frac{\sigma_x^2}{n} = 0 \quad (6)$$

کمترین وریانس و کارایی

چنانچه برآوردگرهای مختلفی برای برآورد یک پارامتر وجود داشته باشند و دارای وریانسهای مختلف باشند، از بین آنها برآوردگری مناسب تر است که کمترین وریانس را دارا باشد. بدیهی است ویژگی ناریبی از اهمیت بالاتری نسبت به کمترین وریانس برخوردار می باشد. بنابراین همواره تلاش می شود از بین برآوردگرهای ناریب، برآوردگری که دارای کمترین وریانس است انتخاب شود. در این صورت کارایی برآوردگر به حداکثر می رسد و به آن برآوردگر کارآ می گویند.

روشهای برآورد

همانطور که اشاره شد روشهای متفاوتی برای برآوردی کمیت وجود دارد ولی بهترین برآورد مربوط به بهترین برآوردگر است که دارای ویژگیهای فوق باشد. در مسائل عملی با استفاده از نمونه ها ممکن است برآوردگرها تمام ویژگیهای فوق را در حالت ایده آل نداشته باشند ولی با بزرگ شدن نمونه ها به این ویژگی ها نزدیک شوند. منظور از برآوردگرها در این بخش برآوردگرهای نقطه ای هستند به برخی از آنها اشاره می شود.

روش گشتاورها: در این روش از چند گشتاور اول یک نمونه متناظر با گشتاورهای جامعه برای به دست آوردن مجهولات مورد نظر خود استفاده می شود. به طور مثال گشتاور مرتبه r برای یک نمونه به صورت زیر است که میانگین با همین وسیله به دست می آید.

$$m_r = \frac{1}{n} \sum_{i=1}^n x_i^r \quad (7)$$

روش حداکثر درست نمایی: این روش برای برآورد پارامترهای مجهول $\underline{\theta} = (\theta_1, \theta_2, \theta_3, \dots, \theta_m)$ به گونه ای بکار می رود که احتمال نمونه به ازای آن حداکثر شود. چنانچه متغیر تصادفی x_i مستقل و دارای تابع چگالی $f(x_i)$ یکسان باشند، در آن صورت بردار نمونه $(x_1, x_2, x_3, \dots, x_n)$ دارای تابع چگالی توام زیر است.

$$F(x_1, x_2, x_3, \dots, x_n; \theta_1, \theta_2, \theta_3, \dots, \theta_m) = \prod_{i=1}^n f(x_i; \theta_1, \theta_2, \theta_3, \dots, \theta_m) \quad (7)$$

تابع فوق تابع درست نمایی نامیده می شود و توابع چگالی $f(x_i)$ توابعی از پارامترهای مجهول $(\theta_1, \theta_2, \theta_3, \dots, \theta_m)$ می باشند. پارامترهای مجهول را می توان با روابط $\theta_i = g(x_1, x_2, x_3, \dots, x_n)$ برآورد نمود. در این روش برآوردهای $\hat{\theta}_i$ باید به گونه ای باشند که رابطه (7) یعنی تابع درست نمایی حداکثر شود. بنابراین مشتق زیر را برقرار می کنیم.

$$\frac{\partial F}{\partial \underline{\theta}} = 0 \quad \text{یا} \quad \frac{\partial \ln(F)}{\partial \underline{\theta}} = 0 \quad (8)$$

حل معادله فوق منجر به برآورد پارامترهای مجهول جامعه با شرط حداکثر احتمال مقادیر نمونه می شود.

روش کمترین مربعات: بر خلاف روش حداکثر درست نمایی در روش کمترین مربعات نیازی به دانستن توابع چگالی اندازه گیری ها نیست. هر چند برای فواصل اطمینان و آزمون های آماری به توابع چگالی اندازه گیری ها نیاز می باشد. در صورتیکه متغیرهای تصادفی دارای دارای تابع توزیع نرمال باشند، نتایج حاصل از روش کمترین مربعات و روش حداکثر درست نمایی یکسان خواهند بود. اساس روش کمترین مربعات بر کمینه سازی مجموع مربعات یا قیمانده های اندازه گیری ها است. اگر بردار باقیمانده $\underline{v} = \underline{\hat{x}} - \underline{x}$ به عنوان

بردار متغیر تصادفی با تابع توزیع نرمال و ویژگی $E(\underline{v}) = \underline{0}$ در نظر گرفته شود، در این صورت تابع چگالی توام باقیمانده ها به صورت زیر نمایش داده می شود.

$$\begin{aligned} f(\underline{v}) &= \frac{1}{(2\pi)^{n/2} \cdot |C_v|^{1/2}} \times \exp\left[-\frac{1}{2}(\underline{v} - E(\underline{v}))^T \underline{C}_v^{-1}(\underline{v} - E(\underline{v}))\right] \\ &= \frac{1}{(2\pi)^{n/2} \cdot |C_v|^{1/2}} \times \exp\left[-\frac{1}{2}\underline{v}^T \underline{C}_v^{-1} \underline{v}\right] \end{aligned} \quad (9)$$

حال با اعمال شرط به حداقل رساندن $\underline{v}^T \underline{C}_v^{-1} \underline{v}$ به عنوان روش کمترین مربعات به یک برآورد از بردار $(x_1, x_2, x_3, \dots, x_n)$ خواهیم رسید. با نگاهی دقیق به رابطه (۹) به روشنی پیداست که حداقل کردن $\underline{v}^T \underline{C}_v^{-1} \underline{v}$ معادل حداکثر سازی $f(\underline{v})$ است که همان روش حداکثر درست نمایی است.

تعیین فواصل اطمینان (برآورد فاصله ای)

فاصله اطمینان برای هر پارامتر θ مربوط به یک جامعه که $\hat{\theta}$ مقدار برآورد شده آن بر اساس یک نمونه است، به صورت زیر تعریف می گردد.

$$P(\theta_1 \leq \theta \leq \theta_2) = 1 - \alpha \quad (10)$$

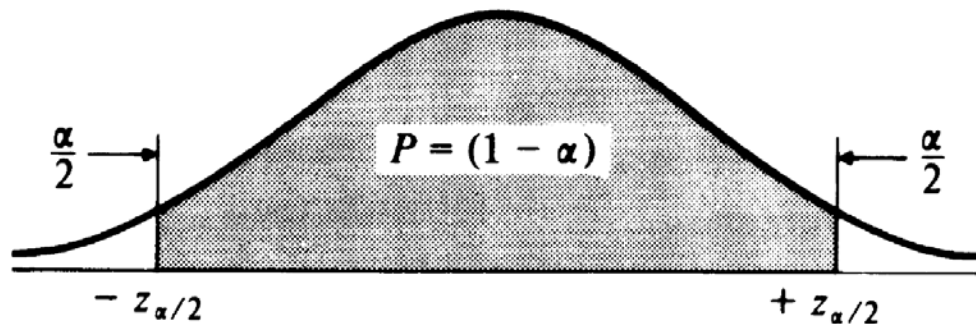
که در آن $1 - \alpha$ سطح اطمینان و α سطح معنی دار نامیده می شود. در واقع فاصله اطمینان $[\theta_1, \theta_2]$ فاصله ای است که انتظار داریم پارامتر θ با احتمال $1 - \alpha$ در آن قرار داشته باشد یا به عبارت دیگر انتظار داریم پارامتر θ با احتمال α در آن قرار نداشته باشد. در مسائل عملی سطوح اطمینان معمولاً $1 - \alpha = 0.95$ یا $1 - \alpha = 0.99$ ($\alpha = 0.05$) یا $1 - \alpha = 0.99$ ($\alpha = 0.01$) در نظر گرفته می شوند. برای یافتن حدود θ_1 و θ_2 یا به عبارت بهتر تعیین فاصله اطمینان $[\theta_1, \theta_2]$ می بایست تابع توزیع برآوردگر مورد نظر را داشته باشیم.

برای مثال می توانیم فاصله اطمینان برآورد میانگین $\bar{x}$ یک جامعه با میانگین و انحراف معیار واقعی μ_x و σ_x را برای سطح معنی دار α به صورت زیر نمایش دهیم. در این مثال تعداد عناصر نمونه را n و آنها دارای تابع توزیع نرمال فرض می کنیم.

$$P\left(\frac{(\bar{x} - z_{\alpha/2})\sigma_x}{\sqrt{n}} \leq \mu_x \leq \frac{(\bar{x} + z_{\alpha/2})\sigma_x}{\sqrt{n}}\right) = 1 - \alpha \quad (11)$$

با یادآوری تعریف متغیر تصادفی نرمال استاندارد، رابطه (۱۱) را می توان به صورت زیر بازنویسی نمود (نگاره ۲).

$$P\left(-z_{\alpha/2} \leq \frac{\bar{x} - \mu_x}{\sigma_x/\sqrt{n}} \leq z_{\alpha/2}\right) = 1 - \alpha \quad (12)$$



نگاره ۲- فاصله اطمینان برای میانگین یک نمونه

میانگین وزن دار

یکی از برآوردگرهای مهم که همه ویژگی های گفته شده برای یک برآوردگر خوب را داراست و کاربرد نسبتاً زیادی در مسایل مهندسی دارد، برآوردگر میانگین وزن دار است که در این بخش به آن می پردازیم. فرض کنید m نمونه اندازه گیری از یک موضوع با حجم های $n_1, n_2, n_3, \dots, n_m$ طی m بار به دست آمده اند.

چنانچه انحراف معیار هر اندازه گیری را مساوی و برابر با σ در نظر گرفته شود، میانگین و وریانس میانگین هر نمونه n_i تایی به صورت زیر به دست می آیند.

$$(۱۳) \quad \begin{cases} \bar{x}_1 = \frac{1}{n_1} \sum_{i=1}^{n_1} x_i, & \sigma_1^2 = \frac{\sigma^2}{n_1} \\ \bar{x}_2 = \frac{1}{n_2} \sum_{i=1}^{n_2} x_i, & \sigma_2^2 = \frac{\sigma^2}{n_2} \\ \bar{x}_3 = \frac{1}{n_3} \sum_{i=1}^{n_3} x_i, & \sigma_3^2 = \frac{\sigma^2}{n_3} \\ \vdots & \vdots \\ \bar{x}_m = \frac{1}{n_m} \sum_{i=1}^{n_m} x_i, & \sigma_m^2 = \frac{\sigma^2}{n_m} \end{cases}$$

از طرفی با یکپارچه فرض کردن تمام اندازه گیری ها، می توان میانگین کلی و و وریانس آن را محاسبه نمود.

در این صورت حجم کل اندازه گیری ها $N = \sum_{i=1}^m n_i$ و مجموع تمام اندازه گیری ها $\sum_{i=1}^N x_i = \sum_{i=1}^m n_i \bar{x}_i$ خواهد بود و میانگین کلی بنابر تعریف به صورت زیر به دست می آید.

$$(۱۴) \quad \bar{x} = \frac{1}{N} \sum_{i=1}^N x_i = \frac{\sum_{i=1}^m n_i \bar{x}_i}{\sum_{i=1}^m n_i}$$

حال تعداد اندازه گیری های هر نمونه را بر پایه تعریف وریانس میانگین هر نمونه به صورت زیر به دست می آوریم.

$$\begin{cases} n_1 = \frac{\sigma^2}{\sigma_1^2}, \\ n_2 = \frac{\sigma^2}{\sigma_2^2}, \\ n_3 = \frac{\sigma^2}{\sigma_3^2}, \\ \vdots \\ n_m = \frac{\sigma^2}{\sigma_m^2} \end{cases} \quad (15)$$

اكنون با استفاده از روابط فوق و جايگذاري آنها در رابطه ميانگين كلي به رابطه جديد زير مي رسييم.

$$\bar{x} = \frac{\sum_{i=1}^m \left(\frac{\sigma^2}{\sigma_i^2} \right) \bar{x}_i}{\sum_{i=1}^m \left(\frac{\sigma^2}{\sigma_i^2} \right)} \quad (16)$$

با يادآوري تعريف وزن بر حسب وريانس $(p_i = \frac{\sigma^2}{\sigma_i^2})$ ، مجدداً ميانگين كلي به صورت زير بازنويسي مي شود.

$$\bar{x} = \frac{\sum_{i=1}^m p_i \bar{x}_i}{\sum_{i=1}^m p_i} \quad (17)$$

با نگاهی دقيق در رابطه اخير، متوجه مي شويم كه تنها با استفاده از نتيجه هر نمونه كوچك (ميانگين نمونه) و وريانس آن، مي توانيم به نتيجه يكساني براي حالي كه تمام اندازه گيري ها در نظر گرفته شوند، برسيم. رابطه اخير در واقع ميانگين وزندار است و در حالت كلي مي توان آن را براي برآورد ميانگين اندازه گيري هاي مختلف (x_i) با وزن هاي متفاوت (p_i) به صورت زير نمايش داد.

$$\bar{x} = \frac{\sum_{i=1}^m p_i x_i}{\sum_{i=1}^m p_i} \quad (18)$$

به طور مشابه با میانگین معمولی، انتظار می رود دقت میانگین وزندار از دقت تک تک اندازه گیری ها بالاتر باشد. برای بدست آوردن دقت میانگین وزندار با استفاده از قانون انتشار خطاها به صورت زیر عمل می کنیم.

$$\sigma_x^2 = \left(\frac{1}{\sum_{i=1}^m p_i} \right)^2 \left[\sum_{i=1}^m p_i^2 \sigma_i^2 \right] \quad (19)$$

با توجه به تعریف وزن می توانیم عملیات زیر را برای رابطه فوق انجام دهیم.

$$\begin{aligned} \sigma_x^2 &= \left(\frac{1}{\sum_{i=1}^m p_i} \right)^2 \left[\sum_{i=1}^m \frac{\sigma_0^4}{\sigma_i^4} \cdot \sigma_i^2 \right] \\ &= \left(\frac{\sigma_0}{\sum_{i=1}^m p_i} \right)^2 \left[\sum_{i=1}^m \frac{\sigma_0^2}{\sigma_i^2} \right] \\ &= \left(\frac{\sigma_0}{\sum_{i=1}^m p_i} \right)^2 \left[\sum_{i=1}^m p_i \right] \end{aligned} \quad (20)$$

در پایان می توانیم رابطه وریانس و انحراف معیار میانگین وزندار را به صورت زیر بازنویسی نماییم.

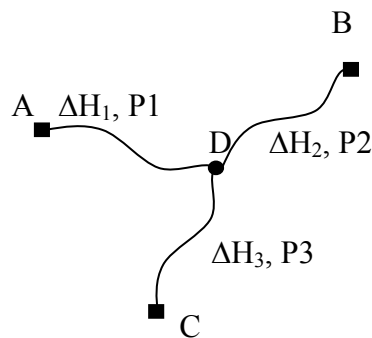
$$\sigma_x^2 = \frac{\sigma_0^2}{\sum_{i=1}^m p_i} \quad (21)$$

$$\sigma_x = \frac{\sigma_0}{\sqrt{\sum_{i=1}^m p_i}} \quad (22)$$

در صورتیکه $\sigma_0 = 1$ در نظر گرفته شود، رابطه انحراف معیار میانگین وزندار به صورت زیر در می آید.

$$\sigma_x = \frac{1}{\sqrt{\sum_{i=1}^m p_i}} \quad (23)$$

برای مثال مطابق نگاره (۳) از سه نقطه ارتفاعی معلوم A ، B و C به نقطه مجهول D ترازیبی شده است. چنانچه اختلاف ارتفاع های ΔH_1 ، ΔH_2 و ΔH_3 به ترتیب با وزن های p_1 ، p_2 و p_3 اندازه گیری شده باشند، محتملترین مقدار ارتفاع نقطه D را به دست آورید.



نگاره ۳- تعیین ارتفاع مجهول D با ترازیبی از از سه نقطه ارتفاعی معلوم A ، B و C

بر اساس رابطه میانگین وزندار، ارتفاع و انحراف معیار نقطه مجهول D به صورت زیر قابل محاسبه است.

$$H_D = \overline{X} = \frac{P_1(\Delta H_1 + H_A) + P_2(H_3 + \Delta H_2) + P_3(\Delta H_3 + H_C)}{P_1 + P_2 + P_3}$$

$$\sigma_{H_D} = \frac{1}{\sqrt{P_1 + P_2 + P_3}}$$

فصل پنجم

نمایش هندسی پخش خطاها

برای بیان توام عدم قطعیت یا خطای چند متغیر تصادفی از کمیت های وریانس و کوریانس بین آنها استفاده می شود. کمیت های فوق عددی می باشند و تصور درستی بویژه هنگامیکه با متغیرهای مکانی سروکار داریم برای ما فراهم نمی سازند. به همین منظور از یک نمایش گرافیکی برای متغیرهای تصادفی استفاده می شود که چگونگی توزیع خطا اطراف نقطه مورد نظر را نشان می دهد. معمولاً این نمایش گرافیکی با بیضی در حالت دو بعدی یا بیضوی در حالت سه بعدی انجام می گیرد.

بیضی خطا

از آنجا که به طور سنتی وابستگی ها بین هر دو متغیر تصادفی ارائه می شوند، لذا ساده ترین حالت ممکن، یعنی در نظر گرفتن متغیر های تصادفی به صورت دو به دو، مورد نظر می باشد. بنابراین در این فصل با نمایش گرافیکی دو بعدی از پخش خطا برای دو متغیر تصادفی سروکار خواهیم داشت. این نمایش گرافیکی یک بیضی است که در واقع بیانگر یک سطح اطمینان از قرار گیری مقادیر واقعی دو متغیر تصادفی به صورت توام در آن است.

برای اثبات اینکه پخش خطا اطراف هر نقطه یک بیضی است، ابتدا تابع چگالی نرمال n متغیره با فرض n متغیر تصادفی $(X_1, X_2, \dots, X_n)$ را که قبلاً به صورت زیر معرفی شده است در نظر بگیرید.

$$f(x_1, x_2, \dots, x_n) = \frac{1}{(2\pi)^{n/2} \cdot |C_x|^{1/2}} \times \exp \left[-\frac{1}{2} (\underline{X} - \underline{\mu}_x)^T \underline{C}_x^{-1} (\underline{X} - \underline{\mu}_x) \right] \quad (1)$$

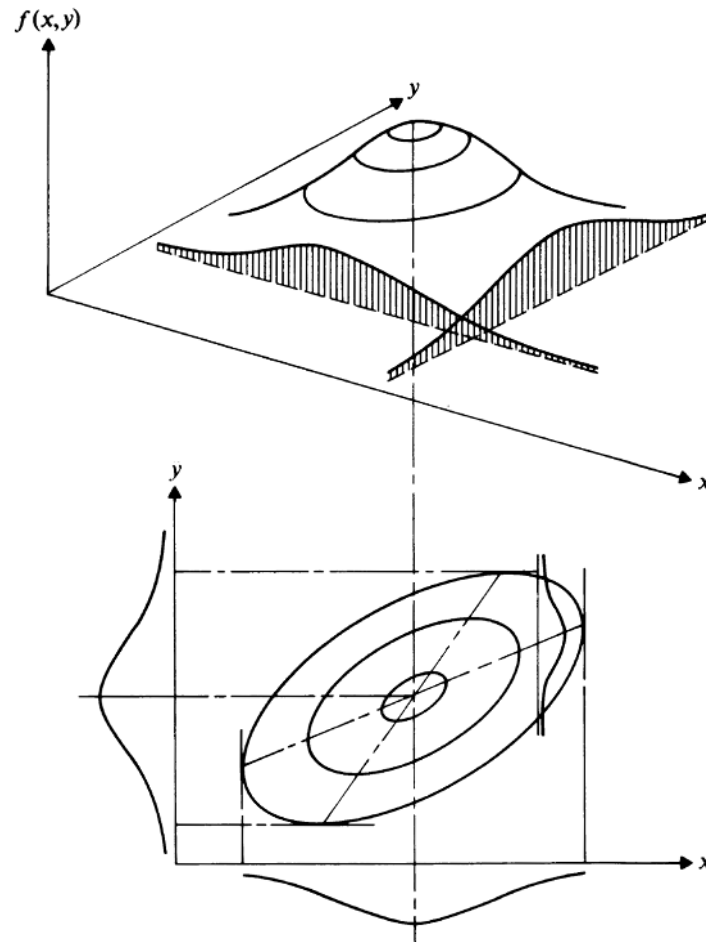
حال با توجه به اینکه فقط دو متغیر تصادفی مورد نظر است، تنها دو متغیر تصادفی x و y در نظر گرفته می شوند و تابع چگالی توام نرمال آنها به صورت زیر خواهد بود (نگاره (۱)).

$$f(x, y) = \frac{1}{2\pi\sqrt{\sigma_x^2\sigma_y^2 - \sigma_{xy}^2}} \exp\left[-\frac{1}{2}(x - \mu_x \quad y - \mu_y)C_x^{-1}(x - \mu_x \quad y - \mu_y)^T\right] \quad (2)$$

$$f(x, y) = \frac{1}{2\pi\sqrt{\sigma_x^2\sigma_y^2 - \sigma_{xy}^2}} \exp\left[-\frac{\sigma_x^2\sigma_y^2}{2(\sigma_x^2\sigma_y^2 - \sigma_{xy}^2)} \times \left(\left(\frac{x - \mu_x}{\sigma_x}\right)^2 - 2\sigma_{xy} \frac{(x - \mu_x)(y - \mu_y)}{\sigma_x^2\sigma_y^2} + \left(\frac{y - \mu_y}{\sigma_y}\right)^2\right)\right] \quad (3)$$

نمایش دیگری از معادله فوق بر حسب ضریب همبستگی بین دو متغیر تصادفی x و y (ρ_{xy}) به صورت زیر است.

$$f(x, y) = \frac{1}{2\pi\sigma_x\sigma_y\sqrt{1 - \rho_{xy}^2}} \exp\left[-\frac{1}{2(1 - \rho_{xy}^2)} \times \left(\left(\frac{x - \mu_x}{\sigma_x}\right)^2 - 2\rho_{xy} \left(\frac{x - \mu_x}{\sigma_x}\right)\left(\frac{y - \mu_y}{\sigma_y}\right) + \left(\frac{y - \mu_y}{\sigma_y}\right)^2\right)\right] \quad (4)$$



نگاره ۱- نمایش تابع چگالی نرمال دو متغیره و تصویر برش مقاطع افقی

چنانچه هر صفحه دلخواه موازی با صفحه xy در یک ارتفاع معین مانند z رویه تابع چگالی نرمال دو متغیره فوق را قطع نماید، مقطع حاصل یک بیضی خواهد بود که به ترتیب زیر اثبات می شود.

$$z = \frac{1}{2\pi\sigma_x\sigma_y\sqrt{1-\rho_{xy}^2}} \exp\left[-\frac{1}{2}K^2\right] \quad (5)$$

حال روابط (۲) و (۵) را با یکدیگر قطع می دهیم و به رابطه زیر می رسیم.

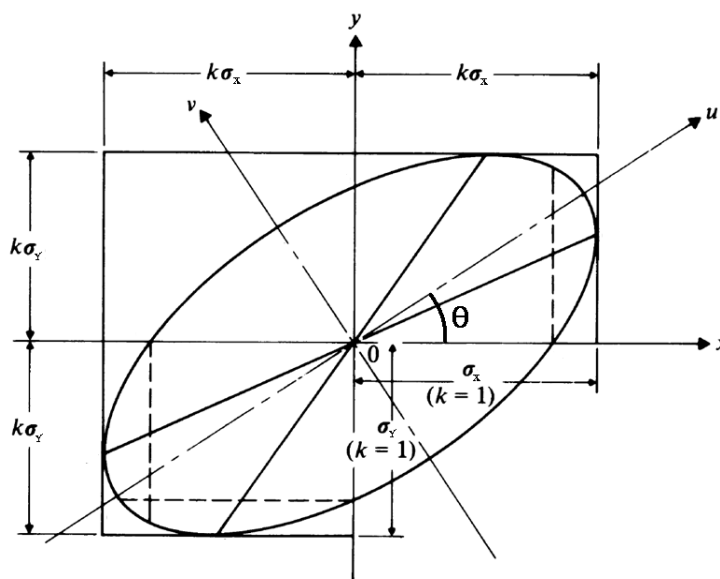
$$\frac{1}{(1-\rho_{xy}^2)} \left[\left(\frac{x-\mu_x}{\sigma_x} \right)^2 - 2\rho_{xy} \left(\frac{x-\mu_x}{\sigma_x} \right) \left(\frac{y-\mu_y}{\sigma_y} \right) + \left(\frac{y-\mu_y}{\sigma_y} \right)^2 \right] = K^2 \quad (6)$$

رابطه (۶) بیانگر یک معادله بیضی در حالت عمومی است که مختصات مرکز آن میانگین های μ_x و μ_y می باشد و در یک مستطیل با ابعاد $2K\sigma_x$ در راستای محور x و $2K\sigma_y$ در راستای محور y ها قرار گرفته است. این بیضی معروف به بیضی خطا است که برای تعیین مشخصات آن به ترتیب زیر عمل می کنیم. ابتدا فرم برداری رابطه (۶) را می نویسیم.

$$(x-\mu_x \quad y-\mu_y) C_x^{-1} (x-\mu_x \quad y-\mu_y)^T = K^2 \quad (7)$$

رابطه (۷) با انتقال به مبدا مختصات به صورت زیر در می آید که بیانگر یک معادله بیضی است که مرکز آن منطبق بر مبدا مختصات می باشد.

$$(x \quad y) C_x^{-1} (x \quad y)^T = K^2 \quad (8)$$



نگاره ۲- بیضی خطا پس از انتقال به مبدا

با دوران محورهای x و y به اندازه θ به منظور انطباق بر محورهای u و v به یک معادله جدید از بیضی خواهیم رسید که قطرهای بزرگ و کوچک آن به ترتیب منطبق بر محورهای u و v خواهد گردید (نگاره (۲)). با انجام مجموعه عملیات زیر به یک ماتریس وریانس کوریانس قطری برای محورهای u و v خواهیم رسید

$$\underline{x} = \begin{bmatrix} x \\ y \end{bmatrix}, \quad \underline{X} = \begin{bmatrix} u \\ v \end{bmatrix}$$

$$\underline{C}_x^{-1} = \begin{bmatrix} \sigma_x^2 & \sigma_{xy} \\ \sigma_{xy} & \sigma_y^2 \end{bmatrix}^{-1} = \frac{1}{\sigma_x^2 \sigma_y^2 - \sigma_{xy}^2} \begin{bmatrix} \sigma_y^2 & -\sigma_{xy} \\ -\sigma_{xy} & \sigma_x^2 \end{bmatrix}$$

$$\underline{R} = \begin{bmatrix} \cos \theta & \sin \theta \\ -\sin \theta & \cos \theta \end{bmatrix}$$

$$\underline{x}^T \underline{C}_x^{-1} \underline{x} = K^2$$

$$\underline{X} = \underline{R} \underline{x}$$

$$\underline{x} = \underline{R}^{-1} \underline{X} = \underline{R}^T \underline{X}$$

$$\underline{X}^T \underline{R} \underline{C}_x^{-1} \underline{R}^T \underline{X} = \underline{X}^T \begin{bmatrix} \sigma_u^2 & 0 \\ 0 & \sigma_v^2 \end{bmatrix}^{-1} \underline{X} = K^2$$

$$\underline{R} \underline{C}_x^{-1} \underline{R}^T = \begin{bmatrix} \sigma_u^2 & 0 \\ 0 & \sigma_v^2 \end{bmatrix}^{-1}$$

$$\underline{RC}_x^{-1} \underline{R}^T = \frac{1}{\sigma_x^2 \sigma_y^2 - \sigma_{xy}^2} \times \begin{bmatrix} \sigma_y^2 \cos^2 \theta - 2\sigma_{xy} \sin \theta \cos \theta + \sigma_x^2 \sin^2 \theta & \sigma_{xy} (\sin^2 \theta - \cos^2 \theta) + \sin \theta \cos \theta (\sigma_x^2 - \sigma_y^2) \\ \sigma_{xy} (\sin^2 \theta - \cos^2 \theta) + \sin \theta \cos \theta (\sigma_x^2 - \sigma_y^2) & \sigma_y^2 \sin^2 \theta + 2\sigma_{xy} \sin \theta \cos \theta + \sigma_x^2 \cos^2 \theta \end{bmatrix}$$

پس از دوران به اندازه θ و قطری سازی ماتریس وریانس کوریانس بردار $\underline{x}$ ، باید عناصر خارج از قطر اصلی صفر شوند. با توجه به چرخش بیضی و قرارگیری اقطار آن بر محورهای u و v ، ساختار ماتریس وریانس وریانس جدید به صورت زیر می باشد.

$$\underline{RC}_x^{-1} \underline{R}^T = \begin{bmatrix} \frac{1}{a^2} & 0 \\ 0 & \frac{1}{b^2} \end{bmatrix}$$

با صفر قرار دادن هریک از عناصر خارج از قطر اصلی در ماتریس $\underline{RC}_x^{-1} \underline{R}^T$ ، می توان مقدار زاویه دوران θ را به عنوان توجیه بیضی خطا پیدا کرد.

$$\sigma_{xy} (\sin^2 \theta - \cos^2 \theta) + \sin \theta \cos \theta (\sigma_x^2 - \sigma_y^2) = -\sigma_{xy} \cos 2\theta + \frac{1}{2} \sin 2\theta (\sigma_x^2 - \sigma_y^2) = 0$$

$$\tan 2\theta = \frac{2\sigma_{xy}}{\sigma_x^2 - \sigma_y^2} \quad (9)$$

چون در ادامه به مقادیر $\sin 2\theta$ و $\cos 2\theta$ احتیاج داریم، آنها را نیز به صورت زیر بدست می آوریم.

$$\frac{\sin 2\theta}{\cos 2\theta} = \frac{2\sigma_{xy}}{\sigma_x^2 - \sigma_y^2} \quad (10)$$

طرفین را به توان ۲ رسانیده و سپس با ترکیب در مخرج به رابطه زیر خواهیم رسید.

$$\frac{\sin^2 2\theta}{1 - \sin^2 2\theta} = \frac{4\sigma_{xy}^2}{(\sigma_x^2 - \sigma_y^2)^2}$$

$$\sin^2 2\theta = \frac{4\sigma_{xy}^2}{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2}$$

$$\sin 2\theta = \frac{2\sigma_{xy}}{\sqrt{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2}} \quad (11)$$

مجدداً طرفین رابطه (۱۰) را به توان ۲ رسانیده و با ترکیب در صورت به رابطه زیر خواهیم رسید.

$$\frac{1 - \sin^2 2\theta}{\cos^2 2\theta} = \frac{4\sigma_{xy}^2}{(\sigma_x^2 - \sigma_y^2)^2}$$

$$\frac{1}{\cos^2 2\theta} = \frac{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2}{(\sigma_x^2 - \sigma_y^2)^2}$$

$$\cos 2\theta = \frac{\sigma_x^2 - \sigma_y^2}{\sqrt{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2}} \quad (12)$$

حال با در نظر گرفتن $\sin^2 \theta = \frac{1 - \cos 2\theta}{2}$ و $\cos^2 \theta = \frac{1 + \cos 2\theta}{2}$ به محاسبه $\sigma_{\min}$ و $\sigma_{\max}$ می

پردازیم.

$$\begin{aligned}
 \sigma_{\min}^2 &= b^2 = \sigma_y^2 \cos^2 \theta - 2\sigma_{xy} \sin \theta \cos \theta + \sigma_x^2 \sin^2 \theta \quad (۱۳) \\
 &= \frac{1}{2} (\sigma_y^2 + \sigma_y^2 \cos 2\theta - 2\sigma_{xy} \sin 2\theta + \sigma_x^2 - \sigma_x^2 \cos 2\theta) \\
 &= \frac{1}{2} (\sigma_y^2 + \sigma_x^2 + \cos 2\theta (\sigma_y^2 - \sigma_x^2) - 2\sigma_{xy} \sin 2\theta) \\
 &= \frac{1}{2} \left(\sigma_y^2 + \sigma_x^2 + \frac{(\sigma_y^2 - \sigma_x^2)(\sigma_x^2 - \sigma_y^2)}{\sqrt{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2}} - \frac{4\sigma_{xy}^2}{\sqrt{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2}} \right) \\
 &= \frac{1}{2} \left(\sigma_y^2 + \sigma_x^2 - \frac{(\sigma_x^2 - \sigma_y^2)^2}{\sqrt{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2}} - \frac{4\sigma_{xy}^2}{\sqrt{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2}} \right) \\
 &= \frac{1}{2} \left(\sigma_y^2 + \sigma_x^2 - \sqrt{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2} \right)
 \end{aligned}$$

$$\begin{aligned}
 \sigma_{\max}^2 &= a^2 = \sigma_y^2 \sin^2 \theta + 2\sigma_{xy} \sin \theta \cos \theta + \sigma_x^2 \cos^2 \theta \quad (۱۴) \\
 &= \frac{1}{2} (\sigma_y^2 - \sigma_y^2 \cos 2\theta + 2\sigma_{xy} \sin 2\theta + \sigma_x^2 + \sigma_x^2 \cos 2\theta) \\
 &= \frac{1}{2} (\sigma_x^2 + \sigma_y^2 + \cos 2\theta (\sigma_x^2 - \sigma_y^2) + 2\sigma_{xy} \sin 2\theta) \\
 &= \frac{1}{2} \left(\sigma_y^2 + \sigma_x^2 + \sqrt{(\sigma_x^2 - \sigma_y^2)^2 + 4\sigma_{xy}^2} \right)
 \end{aligned}$$

$$\frac{1}{\sigma_x^2 \sigma_y^2 - \sigma_{xy}^2} = \frac{1}{a^2 b^2}$$

$$\underline{RC}_x^{-1} \underline{R}^T = \frac{1}{a^2 b^2} \begin{bmatrix} b^2 & 0 \\ 0 & a^2 \end{bmatrix} = \begin{bmatrix} \frac{1}{a^2} & 0 \\ 0 & \frac{1}{b^2} \end{bmatrix}$$

$$\underline{X}^T (\underline{RC}_x^{-1} \underline{R}^T) \underline{X} = \underline{X}^T \begin{bmatrix} a^2 & 0 \\ 0 & b^2 \end{bmatrix}^{-1} \underline{X} = K^2$$

معادله برداری فوق را می توان به صورت زیر که بیانگر یک بیضی با اقطار منطبق بر محورهای u و v است نمایش داد.

$$\frac{u^2}{a^2} + \frac{v^2}{b^2} = K^2 \quad (۱۵)$$

اگر $K = 1$ در نظر گرفته شود در آن صورت به معادله بیضی استاندارد خواهیم رسید.

$$\frac{u^2}{a^2} + \frac{v^2}{b^2} = 1 \quad (۱۶)$$

همانطور که از نگاره (۲) پیداست $a^2 = \sigma_u^2 = \sigma_u^2$ و $b^2 = \sigma_v^2 = \sigma_v^2$ و بنابراین رابطه (۱۶) را به صورت زیر بازنویسی نماییم.

$$\frac{u^2}{\sigma_u^2} + \frac{v^2}{\sigma_v^2} = 1 \quad (۱۷)$$

اگر σ_u^2 و σ_v^2 را در K ضرب کنیم، ابعاد بیضی تغییر خواهد کرد و هر چه قدر K بزرگتر انتخاب شود سطح بیضی خطا بزرگتر و در نتیجه احتمال وجود جواب در داخل بیضی بیشتر خواهد بود. برای تعیین مقدار سطح داخل بیضی باید از تابع توزیع کای اسکور با درجه آزادی ۲ استفاده نماییم. دلیل انتخاب این تابع توزیع این است که اگر دو متغیر تصادفی x و y از تابع توزیع نرمال با میانگین های μ_x و μ_y و واریانس های σ_x^2 و σ_y^2 تبعیت کنند، در آن صورت متغیر کای اسکور حاصل از آنها به صورت زیر می باشد.

$$y = \frac{(x - \mu_x)^2}{\sigma_x^2} + \frac{(y - \mu_y)^2}{\sigma_y^2} \sim \chi_2^2$$

با نگاهی دقیق به معادله بیضی خطا در می یابیم که بیانگر یک متغیر تصادفی کای اسکور با درجه آزادی دو است. لذا به ازای K های مختلف می توان سطوح اطمینان متفاوت به دست آورد (جدول (۱)).

$$P\left(\left[\frac{u^2}{\sigma_u^2} + \frac{v^2}{\sigma_v^2}\right] \leq K^2\right) = P(\chi_{2;\alpha}^2 \leq K^2) = 1 - \alpha \quad (17)$$

جدول ۱- ارتباط سطوح اطمینان و ضریب K در بیضی خطا

K	$1 - \alpha$
1	0.394
1.177	0.50
2.146	0.90
2.447	0.95
3.035	0.99

روشی دیگر برای تعیین عناصر بیضی خطا استفاده از معادله مشخه ماتریس وریانس کوریانس بردار $\underline{x}$ ($|\underline{C}_x - \lambda \underline{I}| = 0$) است. در واقع در این روش نیز با به دست آوردن مقادیر ویژه ماتریس فوق آن را به یک ماتریس قطری تبدیل می کنیم.

$$(-\lambda)^2 + tr(\underline{C}_x)(\lambda) + |\underline{C}_x| = 0 \quad (18)$$

که در آن

$$tr(\underline{C}_x) = \sigma_x^2 + \sigma_y^2$$

$$|\underline{C}_x| = \sigma_x^2 \sigma_y^2 - \sigma_{xy}^2$$

بنابراین معادله مشخصه (۱۸) به صورت زیر قابل بازنویسی است.

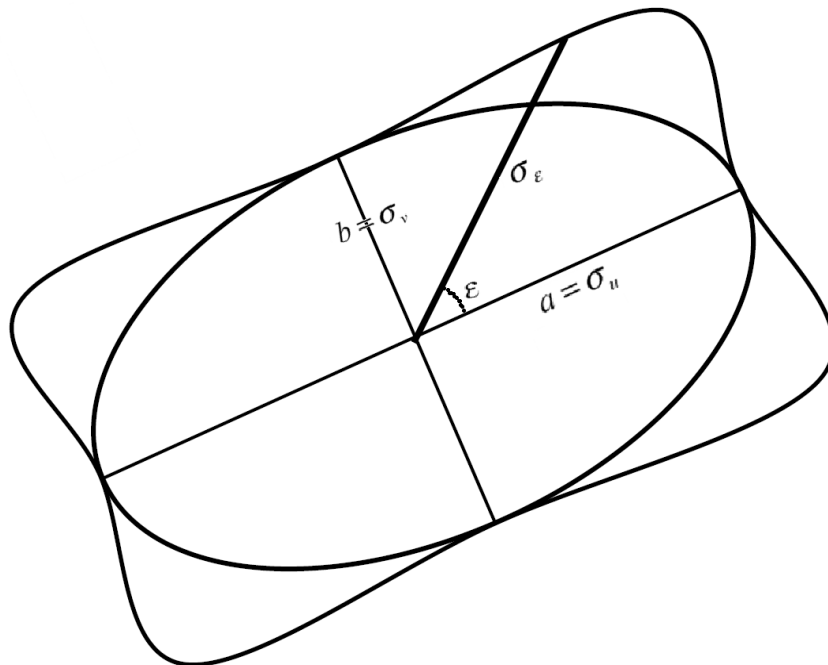
$$\lambda^2 - (\sigma_x^2 + \sigma_y^2)\lambda + (\sigma_x^2 \sigma_y^2 - \sigma_{xy}^2) = 0 \quad (19)$$

جواب معادله فوق $\lambda_1 = a^2$ و $\lambda_2 = b^2$ است ه قبلا نیز به دست آمده بودند.

منحنی پدال

بیضی خطا تنها در دو امتداد u و v خطای استاندارد را به ترتیب به اندازه σ_u^2 و σ_v^2 به دست می دهد و در سایر امتدادها فاصله مرکز بیضی تا منحنی بیضی نشاندهنده خطای استاندارد نمی باشد. چنانچه مایل به محاسبه خطای استاندارد در هر امتداد دلخواه باشیم، رابطه ریاضی بدست آوردن آن برای یک امتداد که با محور بزرگ بیضی خطا زاویه ε در جهت عقربه های ساعت می سازد به صورت زیر به دست می آید. توجه نمایید که منظور از خطای استاندارد در هر امتداد استخراج یک برش یک بعدی از توزیع خطا و در نتیجه تعیین انحراف معیار با احتمال 0.68 است. اثبات رابطه زیر به عهده دانشجو گذاشته می شود.

$$\sigma_\varepsilon^2 = \sigma_u^2 \cos^2 \varepsilon + \sigma_v^2 \sin^2 \varepsilon \quad (20)$$



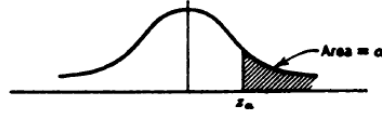
نگاره ۳- منحنی پدال و بیضی خطا

پیوست ها:

جداول آماری

جدول آماری ۱- سطح زیر منحنی تابع چگالی نرمال استاندارد

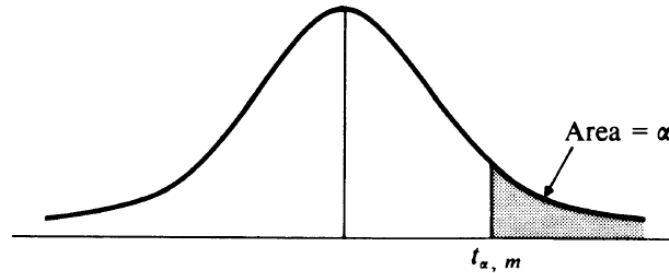
$$\alpha = \int_{z_\alpha}^{\infty} f(z) dz = P(z > z_\alpha) = 1 - \int_{-\infty}^{z_\alpha} f(z) dz$$



$z_\alpha \rightarrow$ ↓	0.00	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
0.0	0.5000	0.4960	0.4920	0.4880	0.4840	0.4801	0.4761	0.4721	0.4681	0.4641
0.1	0.4602	0.4562	0.4522	0.4483	0.4443	0.4404	0.4364	0.4325	0.4286	0.4247
0.2	0.4207	0.4168	0.4129	0.4090	0.4052	0.4013	0.3974	0.3936	0.3897	0.3859
0.3	0.3821	0.3783	0.3745	0.3707	0.3669	0.3632	0.3594	0.3557	0.3520	0.3483
0.4	0.3446	0.3409	0.3372	0.3336	0.3300	0.3264	0.3228	0.3192	0.3156	0.3121
0.5	0.3085	0.3050	0.3015	0.2981	0.2946	0.2912	0.2877	0.2843	0.2810	0.2776
0.6	0.2743	0.2709	0.2676	0.2643	0.2611	0.2578	0.2546	0.2514	0.2483	0.2451
0.7	0.2420	0.2389	0.2358	0.2327	0.2296	0.2266	0.2236	0.2206	0.2177	0.2148
0.8	0.2119	0.2090	0.2061	0.2033	0.2005	0.1977	0.1949	0.1922	0.1894	0.1867
0.9	0.1841	0.1814	0.1788	0.1762	0.1736	0.1711	0.1685	0.1660	0.1635	0.1611
1.0	0.1587	0.1562	0.1539	0.1515	0.1492	0.1469	0.1446	0.1423	0.1401	0.1379
1.1	0.1357	0.1335	0.1314	0.1292	0.1271	0.1251	0.1230	0.1210	0.1190	0.1170
1.2	0.1151	0.1131	0.1112	0.1093	0.1075	0.1056	0.1038	0.1020	0.1003	0.0985
1.3	0.0968	0.0951	0.0934	0.0918	0.0901	0.0885	0.0869	0.0853	0.0838	0.0823
1.4	0.0808	0.0793	0.0778	0.0764	0.0749	0.0735	0.0721	0.0708	0.0694	0.0681
1.5	0.0668	0.0655	0.0643	0.0630	0.0618	0.0606	0.0594	0.0582	0.0571	0.0559
1.6	0.0548	0.0537	0.0526	0.0516	0.0505	0.0495	0.0485	0.0475	0.0465	0.0455
1.7	0.0446	0.0436	0.0427	0.0418	0.0409	0.0401	0.0392	0.0384	0.0375	0.0367
1.8	0.0359	0.0351	0.0344	0.0336	0.0329	0.0322	0.0314	0.0307	0.0301	0.0294
1.9	0.0287	0.0281	0.0274	0.0268	0.0262	0.0256	0.0250	0.0244	0.0239	0.0233
2.0	0.0228	0.0222	0.0217	0.0212	0.0207	0.0202	0.0197	0.0192	0.0188	0.0183
2.1	0.0179	0.0174	0.0170	0.0166	0.0162	0.0158	0.0154	0.0150	0.0146	0.0143
2.2	0.0139	0.0136	0.0132	0.0129	0.0125	0.0122	0.0119	0.0116	0.0113	0.0110
2.3	0.0107	0.0104	0.0102	0.0099	0.0096	0.0094	0.0091	0.0089	0.0087	0.0084
2.4	0.0082	0.0080	0.0078	0.0076	0.0073	0.0071	0.0070	0.0068	0.0066	0.0064
2.5	0.0062	0.0060	0.0059	0.0057	0.0055	0.0054	0.0052	0.0051	0.0049	0.0048
2.6	0.0047	0.0045	0.0044	0.0043	0.0042	0.0040	0.0039	0.0038	0.0037	0.0036
2.7	0.0035	0.0034	0.0033	0.0032	0.0031	0.0030	0.0029	0.0028	0.0027	0.0026
2.8	0.0026	0.0025	0.0024	0.0023	0.0023	0.0022	0.0021	0.0021	0.0020	0.0019
2.9	0.0019	0.0018	0.0018	0.0017	0.0016	0.0016	0.0015	0.0015	0.0014	0.0014
3.0	0.0014	0.0013	0.0013	0.0012	0.0012	0.0011	0.0011	0.0011	0.0010	0.0010
3.1	0.0010	0.0009	0.0009	0.0009	0.0008	0.0008	0.0008	0.0008	0.0007	0.0007
3.2	0.0007	0.0007	0.0006	0.0006	0.0006	0.0006	0.0006	0.0005	0.0005	0.0005
3.3	0.0005	0.0005	0.0005	0.0004	0.0004	0.0004	0.0004	0.0004	0.0004	0.0004
3.4	0.0003	0.0003	0.0003	0.0003	0.0003	0.0003	0.0003	0.0003	0.0003	0.0002

جدول آماری ۲- سطح زیر منحنی تابع چگالی t-student

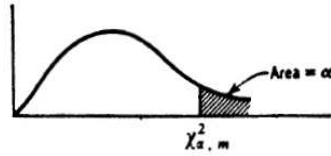
$$t_{\alpha, m} \text{ such that } P(t_m > t_{\alpha, m}) = \alpha = \int_{t_{\alpha, m}}^{\infty} f(t) dt = 1 - \int_{-\infty}^{t_{\alpha, m}} f(t) dt$$



m	$\alpha = 0.25$	0.20	0.15	0.10	0.050	0.025	0.010	0.005	0.0005
1	1.001	1.376	1.963	3.078	6.314	12.706	31.821	63.657	636.619
2	0.816	1.061	1.386	1.886	2.920	4.303	6.965	9.925	31.598
3	0.765	0.978	1.250	1.638	2.353	3.182	4.541	5.841	12.941
4	0.741	0.941	1.190	1.533	2.132	2.776	3.747	4.604	8.610
5	0.727	0.920	1.156	1.476	2.015	2.571	3.365	4.032	6.859
6	0.718	0.906	1.134	1.440	1.943	2.447	3.143	3.707	5.959
7	0.711	0.896	1.119	1.415	1.895	2.365	2.998	3.499	5.405
8	0.706	0.889	1.108	1.397	1.860	2.306	2.896	3.355	5.041
9	0.703	0.883	1.100	1.383	1.833	2.262	2.821	3.250	4.781
10	0.700	0.879	1.093	1.372	1.812	2.228	2.764	3.169	4.587
11	0.697	0.876	1.088	1.363	1.796	2.201	2.718	3.106	4.437
12	0.695	0.873	1.083	1.356	1.782	2.179	2.681	3.055	4.318
13	0.694	0.870	1.079	1.350	1.771	2.160	2.650	3.012	4.221
14	0.692	0.868	1.076	1.345	1.761	2.145	2.624	2.977	4.140
15	0.691	0.866	1.074	1.341	1.753	2.131	2.602	2.947	4.073
16	0.690	0.866	1.071	1.337	1.746	2.120	2.583	2.921	4.015
17	0.689	0.863	1.069	1.333	1.740	2.110	2.567	2.898	3.965
18	0.688	0.862	1.067	1.330	1.734	2.101	2.552	2.878	3.922
19	0.688	0.861	1.066	1.328	1.729	2.093	2.539	2.861	3.883
20	0.687	0.860	1.064	1.325	1.725	2.086	2.528	2.845	3.850
21	0.686	0.859	1.063	1.323	1.721	2.080	2.518	2.831	3.819
22	0.686	0.858	1.061	1.321	1.717	2.074	2.508	2.819	3.792
23	0.685	0.858	1.060	1.319	1.714	2.069	2.500	2.807	3.762
24	0.685	0.857	1.059	1.318	1.711	2.064	2.492	2.797	3.745
25	0.684	0.856	1.058	1.316	1.708	2.060	2.485	2.787	3.725
26	0.684	0.856	1.058	1.315	1.706	2.056	2.479	2.779	3.707
27	0.684	0.855	1.057	1.314	1.703	2.052	2.473	2.771	3.690
28	0.683	0.855	1.056	1.313	1.701	2.048	2.467	2.763	3.674
29	0.683	0.854	1.055	1.311	1.699	2.045	2.462	2.756	3.659
30	0.683	0.854	1.055	1.310	1.697	2.042	2.457	2.750	3.646
40	0.681	0.851	1.050	1.303	1.684	2.021	2.423	2.704	3.551
60	0.679	0.848	1.046	1.296	1.671	2.000	2.390	2.660	3.460
120	0.677	0.844	1.042	1.289	1.658	1.980	2.358	2.617	3.380

جدول آماری ۳- سطح زیر منحنی تابع چگالی کای اسکور

$$\chi^2_{\alpha, m} \text{ such that } P(\chi^2_m > \chi^2_{\alpha, m}) = \alpha = \int_{\chi^2_{\alpha, m}}^{\infty} f(\chi^2) d\chi^2 = 1 - \int_0^{\chi^2_{\alpha, m}} f(\chi^2) d\chi^2$$

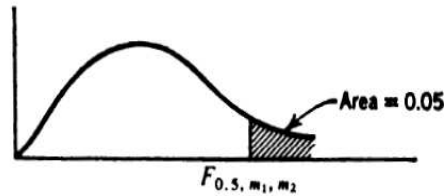


m	$\alpha = 0.995$	0.990	0.975	0.950	0.900	0.500	0.10	0.05	0.025	0.010	0.005
1	0.00	0.00	0.00	0.00	0.02	0.46	2.71	3.84	5.02	6.63	7.88
2	0.01	0.02	0.05	0.10	0.21	1.39	4.61	5.99	7.38	9.21	10.60
3	0.07	0.12	0.22	0.35	0.58	2.37	6.25	7.81	9.35	11.34	12.84
4	0.21	0.30	0.48	0.71	1.06	3.36	7.78	9.49	11.14	13.28	14.86
5	0.41	0.55	0.83	1.15	1.61	4.35	9.24	11.07	12.83	15.09	16.75
6	0.68	0.87	1.24	1.64	2.20	5.35	10.64	12.59	14.45	16.81	18.55
7	0.99	1.24	1.69	2.17	2.83	6.35	12.02	14.07	16.01	18.48	20.28
8	1.34	1.65	2.18	2.73	3.49	7.34	13.36	15.51	17.53	20.09	21.96
9	1.73	2.09	2.70	3.33	4.17	8.34	14.68	16.92	19.02	21.67	23.59
10	2.16	2.56	3.25	3.94	4.87	9.34	15.99	18.31	20.48	23.21	25.19
11	2.60	3.05	3.82	4.57	5.58	10.34	17.28	19.68	21.92	24.73	26.76
12	3.07	3.57	4.40	5.23	6.30	11.34	18.55	21.03	23.34	26.22	28.30
13	3.57	4.11	5.01	5.89	7.04	12.34	19.81	22.36	24.74	27.69	29.82
14	4.07	4.66	5.63	6.57	7.79	13.34	21.06	23.68	26.12	29.14	31.32
15	4.60	5.23	6.26	7.26	8.55	14.34	22.31	25.00	27.49	30.58	32.80
16	5.14	5.81	6.91	7.96	9.31	15.34	23.54	26.30	28.85	32.00	34.27
17	5.70	6.41	7.56	8.67	10.08	16.34	24.77	27.59	30.19	33.41	35.72
18	6.26	7.01	8.23	9.39	10.86	17.34	25.99	28.87	31.53	34.81	37.16
19	6.84	7.63	8.91	10.12	11.65	18.34	27.20	30.14	32.85	36.19	38.58
20	7.43	8.26	9.59	10.85	12.44	19.34	28.41	31.41	34.17	37.57	40.00
21	8.03	8.90	10.28	11.59	13.24	20.34	29.62	32.67	35.48	38.93	41.40
22	8.64	9.54	10.98	12.34	14.04	21.34	30.81	33.92	36.78	40.29	42.80
23	9.26	10.20	11.69	13.09	14.85	22.34	32.01	35.17	38.08	41.64	44.18
24	9.89	10.86	12.40	13.85	15.66	23.34	33.20	36.42	39.36	42.98	45.56
25	10.52	11.52	13.12	14.61	16.47	24.34	34.38	37.65	40.65	44.31	46.93
26	11.16	12.20	13.84	15.38	17.29	25.34	35.56	38.88	41.92	45.64	49.29
27	11.81	12.88	14.57	16.15	18.11	26.34	36.74	40.11	43.19	46.96	49.64
28	12.46	13.56	15.31	16.93	18.94	27.34	37.92	41.34	44.46	48.28	50.99
29	13.12	14.26	16.05	17.71	19.77	28.34	39.09	42.56	45.72	49.59	52.34
30	13.79	14.95	16.79	18.49	20.60	29.34	40.26	43.77	46.98	50.89	53.67
40	20.71	22.16	24.43	26.51	29.05	39.34	51.81	55.76	59.34	63.69	66.77
60	35.53	37.48	40.48	43.19	46.64	59.33	74.40	79.08	83.30	88.38	91.95
120	83.85	86.92	91.58	95.70	100.62	119.33	140.23	146.57	152.21	158.95	163.65

جدول آماری ۴- سطح زیر منحنی تابع چگالی F

$$F_{0.05, m_1, m_2} \text{ such that } P(F_{m_1, m_2} > F_{0.05, m_1, m_2}) = 0.05$$

$$= \int_{F_{0.05, m_1, m_2}}^{\infty} f(F) dF = 1 - \int_0^{F_{0.05, m_1, m_2}} f(F) dF$$



$m_1 \backslash m_2$	1	2	3	4	5	6	7	8	9	10	11	12	13	14	16
1	161	200	216	225	230	234	237	239	241	242	243	244	245	245	246
2	18.5	19.0	19.2	19.2	19.3	19.3	19.4	19.4	19.4	19.4	19.4	19.4	19.4	19.4	19.4
3	10.1	9.55	9.28	9.12	9.01	8.94	8.89	8.85	8.81	8.79	8.76	8.74	8.73	8.71	8.69
4	7.71	6.94	6.59	6.39	6.26	6.16	6.09	6.04	6.00	5.96	5.94	5.91	5.89	5.87	5.84
5	6.61	5.79	5.41	5.19	5.05	4.95	4.88	4.82	4.77	4.73	4.70	4.68	4.66	4.64	4.60
6	5.99	5.14	4.76	4.53	4.39	4.28	4.21	4.15	4.10	4.06	4.03	4.00	3.98	3.96	3.92
7	5.59	4.74	4.35	4.12	3.97	3.87	3.79	3.73	3.68	3.64	3.60	3.57	3.55	3.53	3.49
8	5.32	4.46	4.07	3.84	3.69	3.58	3.50	3.44	3.39	3.35	3.31	3.28	3.26	3.24	3.20
9	5.12	4.26	3.86	3.63	3.48	3.37	3.29	3.23	3.18	3.14	3.10	3.07	3.05	3.03	2.99
10	4.96	4.10	3.71	3.48	3.33	3.22	3.14	3.07	3.02	2.98	2.94	2.91	2.89	2.86	2.83
11	4.84	3.98	3.59	3.36	3.20	3.09	3.01	2.95	2.90	2.85	2.82	2.79	2.76	2.74	2.70
12	4.75	3.89	3.49	3.25	3.11	3.00	2.91	2.85	2.80	2.75	2.72	2.69	2.66	2.64	2.60
13	4.67	3.81	3.41	3.18	3.03	2.92	2.83	2.77	2.71	2.67	2.63	2.60	2.58	2.55	2.51
14	4.60	3.74	3.35	3.11	2.96	2.85	2.76	2.70	2.65	2.60	2.57	2.53	2.51	2.48	2.44
16	4.49	3.63	3.24	3.01	2.85	2.74	2.66	2.59	2.54	2.49	2.46	2.42	2.40	2.37	2.33
18	4.41	3.55	3.16	2.93	2.77	2.66	2.58	2.51	2.46	2.41	2.37	2.34	2.31	2.29	2.25
20	4.35	3.49	3.10	2.87	2.71	2.60	2.51	2.45	2.39	2.35	2.31	2.28	2.25	2.22	2.18
22	4.30	3.44	3.05	2.82	2.66	2.55	2.46	2.40	2.34	2.30	2.26	2.23	2.20	2.17	2.13
24	4.26	3.40	3.01	2.78	2.62	2.51	2.42	2.36	2.30	2.25	2.21	2.18	2.15	2.13	2.09
26	4.23	3.37	2.98	2.74	2.59	2.47	2.39	2.32	2.27	2.22	2.18	2.15	2.12	2.09	2.05
28	4.20	3.34	2.95	2.71	2.56	2.45	2.36	2.29	2.24	2.19	2.15	2.12	2.09	2.06	2.02
30	4.17	3.32	2.92	2.69	2.53	2.42	2.33	2.27	2.21	2.16	2.13	2.09	2.06	2.04	1.99
40	4.08	3.23	2.84	2.61	2.45	2.34	2.25	2.18	2.12	2.08	2.04	2.00	1.97	1.95	1.90
50	4.03	3.18	2.79	2.56	2.40	2.29	2.20	2.13	2.07	2.03	1.99	1.95	1.92	1.89	1.85
60	4.00	3.15	2.76	2.53	2.37	2.25	2.17	2.10	2.04	1.99	1.95	1.92	1.89	1.86	1.82
80	3.96	3.11	2.72	2.49	2.33	2.21	2.13	2.06	2.00	1.95	1.91	1.88	1.84	1.82	1.77
100	3.94	3.09	2.70	2.46	2.31	2.19	2.10	2.03	1.97	1.93	1.89	1.85	1.82	1.79	1.75
200	3.89	3.04	2.65	2.42	2.26	2.14	2.06	1.98	1.93	1.88	1.84	1.80	1.77	1.74	1.69
500	3.86	3.01	2.62	2.39	2.23	2.12	2.03	1.96	1.90	1.85	1.81	1.77	1.74	1.71	1.66
∞	3.84	3.00	2.60	2.37	2.21	2.10	2.01	1.94	1.88	1.83	1.79	1.75	1.72	1.69	1.64

جدول آماری ۴- سطح زیر منحنی تابع چگالی F (ادامه)

$$F_{0.05, m_1, m_2} \text{ such that } P(F_{m_1, m_2} > F_{0.05, m_1, m_2}) = 0.05$$

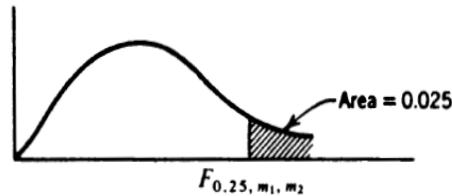
$$= \int_{F_{0.05, m_1, m_2}}^{\infty} f(F) dF = 1 - \int_0^{F_{0.05, m_1, m_2}} f(F) dF$$

18	20	22	24	26	28	30	40	50	60	80	100	200	500	∞	$m_1 \backslash m_2$
247	248	249	249	249	250	250	251	252	252	252	253	254	254	254	1
19.4	19.5	19.5	19.5	19.5	19.5	19.5	19.5	19.5	19.5	19.5	19.5	19.5	19.5	19.5	2
8.67	8.66	8.65	8.64	8.63	8.62	8.62	8.59	8.59	8.57	8.56	8.55	8.54	8.53	8.53	3
5.82	5.80	5.79	5.77	5.76	5.75	5.75	5.72	5.70	5.69	5.67	5.66	5.65	5.64	5.63	4
4.58	3.56	4.54	4.53	4.52	4.50	4.50	4.46	4.44	4.43	4.41	4.41	4.39	4.37	4.37	5
3.90	3.87	3.86	3.84	3.83	3.82	3.81	3.77	3.75	3.74	3.72	3.71	3.69	3.68	3.67	6
3.47	3.44	3.43	3.41	3.40	3.39	3.38	3.34	3.32	3.30	3.29	3.27	3.25	3.24	3.23	7
3.17	3.15	3.13	3.12	3.10	3.09	3.08	3.04	3.02	3.01	2.99	2.97	2.95	2.94	2.93	8
2.96	2.94	2.92	2.90	2.89	2.87	2.86	2.83	2.80	2.79	2.77	2.76	2.73	2.72	2.71	9
2.80	2.77	2.75	2.74	2.72	2.71	2.70	2.66	2.64	2.62	2.60	2.59	2.56	2.55	2.54	10
2.67	2.65	2.63	2.61	2.59	2.58	2.57	2.53	2.51	2.49	2.47	2.46	2.43	2.42	2.40	11
2.57	2.54	2.52	2.51	2.49	2.48	2.47	2.43	2.40	2.38	2.36	2.35	2.32	2.31	2.30	12
2.48	2.46	2.44	2.42	2.41	2.39	2.38	2.34	2.31	2.30	2.27	2.26	2.23	2.22	2.21	13
2.41	2.38	2.37	2.35	2.33	2.32	2.31	2.27	2.24	2.22	2.20	2.19	2.16	2.14	2.13	14
2.30	2.28	2.25	2.24	2.22	2.21	2.19	2.15	2.12	2.11	2.08	2.07	2.04	2.02	2.01	16
2.22	2.19	2.17	2.15	2.13	2.12	2.11	2.06	2.04	2.02	1.99	1.98	1.95	1.93	1.92	18
2.15	2.12	2.10	2.08	2.07	2.05	2.04	1.99	1.97	1.95	1.92	1.91	1.88	1.86	1.84	20
2.10	2.07	2.05	2.03	2.01	2.00	1.98	1.94	1.91	1.89	1.86	1.85	1.82	1.80	1.78	22
2.05	2.03	2.00	1.98	1.97	1.95	1.94	1.89	1.86	1.84	1.82	1.80	1.77	1.75	1.73	24
2.02	1.99	1.97	1.95	1.93	1.91	1.90	1.84	1.82	1.80	1.78	1.76	1.73	1.71	1.69	26
1.99	1.96	1.93	1.91	1.90	1.88	1.87	1.82	1.79	1.77	1.74	1.73	1.69	1.67	1.65	28
1.96	1.93	1.91	1.89	1.87	1.85	1.84	1.79	1.76	1.74	1.71	1.70	1.66	1.64	1.62	30
1.87	1.84	1.81	1.79	1.77	1.76	1.74	1.69	1.66	1.64	1.61	1.59	1.55	1.53	1.51	40
1.81	1.78	1.76	1.74	1.72	1.70	1.69	1.63	1.60	1.58	1.54	1.52	1.48	1.46	1.44	50
1.78	1.75	1.72	1.70	1.68	1.66	1.65	1.59	1.56	1.53	1.50	1.48	1.44	1.41	1.39	60
1.73	1.70	1.68	1.65	1.63	1.62	1.60	1.54	1.51	1.48	1.45	1.43	1.38	1.35	1.32	80
1.71	1.68	1.65	1.63	1.61	1.59	1.57	1.52	1.48	1.45	1.41	1.39	1.34	1.31	1.28	100
1.66	1.62	1.60	1.57	1.55	1.53	1.52	1.46	1.41	1.39	1.35	1.32	1.26	1.22	1.19	200
1.62	1.59	1.56	1.54	1.52	1.50	1.48	1.42	1.38	1.34	1.30	1.28	1.21	1.16	1.11	500
1.60	1.57	1.54	1.52	1.50	1.48	1.46	1.39	1.35	1.32	1.27	1.24	1.17	1.11	1.00	∞

جدول آماری ۴- سطح زیر منحنی تابع چگالی F (ادامه)

$$F_{0.025, m_1, m_2} \text{ such that } P(F_{m_1, m_2} > F_{0.025, m_1, m_2}) = 0.025$$

$$= \int_{F_{0.025, m_1, m_2}}^{\infty} f(F) dF = 1 - \int_0^{F_{0.025, m_1, m_2}} f(F) dF$$



$m_1 \backslash m_2$	1	2	3	4	5	6	7	8	9	10	11	12	13	14	16
1	648	800	864	900	922	937	948	957	963	969	973	977	980	983	987
2	38.5	39.0	39.2	39.2	39.3	39.3	39.4	39.4	39.4	39.4	39.4	39.4	39.4	39.4	39.4
3	17.4	16.0	15.4	15.1	14.9	14.7	14.6	14.5	14.5	14.4	14.4	14.3	14.3	14.3	14.2
4	12.2	10.6	9.98	9.60	9.36	9.20	9.07	8.98	8.90	8.84	8.79	8.75	8.72	8.69	8.64
5	10.0	8.43	7.76	7.39	7.15	6.98	6.85	6.76	6.68	6.62	6.57	6.52	6.49	6.46	6.41
6	8.81	7.26	6.60	6.23	5.99	5.82	5.70	5.60	5.52	5.46	5.41	5.37	5.33	5.30	5.25
7	8.07	6.54	5.89	5.52	5.29	5.12	4.99	4.90	4.82	4.76	4.71	4.67	4.63	4.60	4.54
8	7.57	6.06	5.42	5.05	4.82	4.65	4.53	4.43	4.36	4.30	4.24	4.20	4.16	4.13	4.08
9	7.21	5.71	5.08	4.72	4.48	4.32	4.20	4.10	4.03	3.96	3.91	3.87	3.83	3.80	3.74
10	6.94	5.46	4.83	4.47	4.24	4.07	3.95	3.85	3.78	3.72	3.66	3.62	3.58	3.55	3.50
11	6.72	5.26	4.63	4.28	4.04	3.88	3.76	3.66	3.59	3.53	3.47	3.43	3.39	3.36	3.30
12	6.55	5.10	4.47	4.12	3.89	3.73	3.61	3.51	3.44	3.37	3.32	3.28	3.24	3.21	3.15
13	6.41	4.97	4.35	4.00	3.77	3.60	3.48	3.39	3.31	3.25	3.20	3.15	3.12	3.08	3.03
14	6.30	4.86	4.24	3.89	3.66	3.50	3.38	3.29	3.21	3.15	3.09	3.05	3.01	2.98	2.92
16	6.12	4.69	4.08	3.73	3.50	3.34	3.22	3.12	3.05	2.99	2.93	2.89	2.85	2.82	2.76
18	5.98	4.56	3.95	3.61	3.38	3.22	3.10	3.01	2.93	2.87	2.81	2.77	2.73	2.70	2.64
20	5.87	4.46	3.86	3.51	3.29	3.13	3.01	2.91	2.84	2.77	2.72	2.68	2.64	2.60	2.55
22	5.79	4.38	3.78	3.44	3.22	3.05	2.93	2.84	2.76	2.70	2.65	2.60	2.56	2.53	2.47
24	5.72	4.32	3.72	3.38	3.15	2.99	2.87	2.78	2.70	2.64	2.59	2.54	2.50	2.47	2.41
26	5.66	4.27	3.67	3.33	3.10	2.94	2.82	2.73	2.65	2.59	2.54	2.49	2.45	2.42	2.36
28	5.61	4.22	3.63	3.29	3.06	2.90	2.78	2.69	2.61	2.55	2.49	2.45	2.41	2.37	2.32
30	5.57	4.18	3.59	3.25	3.03	2.87	2.75	2.65	2.57	2.51	2.46	2.41	2.37	2.34	2.28
40	5.42	4.05	3.46	3.13	2.90	2.74	2.62	2.53	2.45	2.39	2.33	2.29	2.25	2.21	2.15
50	5.34	3.98	3.39	3.06	2.83	2.67	2.55	2.46	2.38	2.32	2.26	2.22	2.18	2.14	2.08
60	5.29	3.93	3.34	3.01	2.79	2.63	2.51	2.41	2.33	2.27	2.22	2.17	2.13	2.09	2.03
80	5.22	3.86	3.28	2.95	2.73	2.57	2.45	2.36	2.38	2.21	2.16	2.11	2.07	2.03	1.97
100	5.18	3.83	3.25	2.92	2.70	2.54	2.42	2.32	2.24	2.18	2.12	2.08	2.04	2.00	1.94
200	5.10	3.76	3.18	2.85	2.63	2.47	2.35	2.26	2.18	2.11	2.06	2.01	1.97	1.93	1.87
500	5.05	3.72	3.14	2.81	2.59	2.43	2.31	2.22	2.14	2.07	2.02	1.97	1.93	1.89	1.83
∞	5.02	3.69	3.12	2.79	2.57	2.41	2.29	2.19	2.11	2.05	1.99	1.94	1.90	1.87	1.80

جدول آماری ۴- سطح زیر منحنی تابع چگالی F (ادامه)

$F_{0.025, m_1, m_2}$ such that $P(F_{m_1, m_2} > F_{0.025, m_1, m_2}) = 0.025$

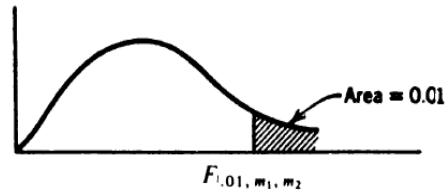
$$= \int_{F_{0.025, m_1, m_2}}^{\infty} f(F) dF = 1 - \int_0^{F_{0.025, m_1, m_2}} f(F) dF$$

18	20	22	24	26	28	30	40	50	60	80	100	200	500	∞	$m_1 \backslash m_2$
990	993	995	997	999	1000	1001	1006	1008	1010	1012	1013	1016	1017	1018	1
39.4	39.4	39.5	39.5	39.5	39.5	39.5	39.5	39.5	39.5	39.5	39.5	39.5	39.5	39.5	2
14.2	14.2	14.1	14.1	14.1	14.1	14.1	14.0	14.0	14.0	14.0	14.0	13.9	13.9	13.9	3
8.60	8.56	8.53	8.51	8.49	8.48	8.46	8.41	8.38	8.36	8.33	8.32	8.29	8.27	8.26	4
6.37	6.33	6.30	6.28	6.26	6.24	6.23	6.18	6.14	6.12	6.10	6.08	6.05	6.03	6.01	5
5.21	5.17	5.14	5.12	5.10	5.08	5.07	5.01	4.98	4.96	4.93	4.92	4.88	4.86	4.85	6
4.50	4.47	4.44	4.42	4.39	4.38	4.36	4.31	4.28	4.25	4.23	4.21	4.18	4.16	4.14	7
4.03	4.00	3.97	3.95	3.93	3.91	3.89	3.84	3.81	3.78	3.76	3.74	3.70	3.68	3.67	8
3.70	3.67	3.64	3.61	3.59	3.58	3.56	3.51	3.47	3.45	3.42	3.40	3.37	3.35	3.33	9
3.45	3.42	3.39	3.37	3.34	3.33	3.31	3.26	3.22	3.20	3.17	3.15	3.12	3.09	3.08	10
3.26	3.23	3.20	3.17	3.15	3.13	3.12	3.06	3.03	3.00	2.97	2.96	2.92	2.90	2.88	11
3.11	3.07	3.04	3.02	3.00	2.98	2.96	2.91	2.87	2.85	2.82	2.80	2.76	2.74	2.72	12
2.98	2.95	2.92	2.89	2.87	2.85	2.84	2.78	2.74	2.72	2.69	2.67	2.63	2.61	2.60	13
2.88	2.84	2.81	2.79	2.77	2.75	2.73	2.67	2.64	2.61	2.58	2.56	2.53	2.50	2.49	14
2.72	2.68	2.65	2.63	2.60	2.58	2.57	2.51	2.47	2.45	2.42	2.40	2.36	2.33	2.32	16
2.60	2.56	2.53	2.50	2.48	2.46	2.44	2.38	2.35	2.32	2.29	2.27	2.23	2.20	2.19	18
2.50	2.46	2.43	2.41	2.39	2.37	2.35	2.29	2.25	2.22	2.19	2.17	2.13	2.10	2.09	20
2.43	2.39	2.36	2.33	2.31	2.29	2.27	2.21	2.17	2.14	2.11	2.09	2.05	2.02	2.00	22
2.36	2.33	2.30	2.27	2.25	2.23	2.21	2.15	2.11	2.08	2.05	2.02	1.98	1.95	1.94	24
2.31	2.28	2.24	2.22	2.19	2.17	2.16	2.09	2.05	2.03	1.99	1.97	1.92	1.90	1.88	26
2.27	2.23	2.20	2.17	2.15	2.13	2.11	2.05	2.01	1.98	1.94	1.92	1.88	1.85	1.83	28
2.23	2.20	2.16	2.14	2.11	2.09	2.07	2.01	1.97	1.94	1.90	1.88	1.84	1.81	1.79	30
2.11	2.07	2.03	2.01	1.98	1.96	1.94	1.88	1.83	1.80	1.76	1.74	1.69	1.66	1.64	40
2.03	1.99	1.96	1.93	1.91	1.88	1.87	1.80	1.75	1.72	1.68	1.66	1.60	1.57	1.55	50
1.98	1.94	1.91	1.88	1.86	1.83	1.82	1.74	1.70	1.67	1.62	1.60	1.54	1.51	1.48	60
1.93	1.88	1.85	1.82	1.79	1.77	1.75	1.68	1.63	1.60	1.55	1.53	1.47	1.43	1.40	80
1.89	1.85	1.81	1.78	1.76	1.74	1.71	1.64	1.59	1.56	1.51	1.48	1.42	1.38	1.35	100
1.82	1.78	1.74	1.71	1.68	1.66	1.64	1.56	1.51	1.47	1.42	1.39	1.32	1.27	1.23	200
1.78	1.74	1.70	1.67	1.64	1.62	1.60	1.51	1.46	1.42	1.37	1.34	1.25	1.19	1.14	500
1.75	1.71	1.67	1.64	1.61	1.59	1.57	1.48	1.43	1.39	1.33	1.30	1.21	1.13	1.00	∞

جدول آماری ۴- سطح زیر منحنی تابع چگالی F (ادامه)

$$F_{0.01, m_1, m_2} \text{ such that } P(F_{m_1, m_2} > F_{0.01, m_1, m_2}) = 0.01$$

$$= \int_{F_{0.01, m_1, m_2}}^{\infty} f(F) dF = 1 - \int_0^{F_{0.01, m_1, m_2}} f(F) dF$$



$m_1 \backslash m_2$	1	2	3	4	5	6	7	8	9	10	11	12	13	14	16
*1	405	500	540	563	576	586	593	598	602	606	608	611	613	614	617
2	98.5	99.0	99.2	99.2	99.3	99.3	99.4	99.4	99.4	99.4	99.4	99.4	99.4	99.4	99.4
3	34.1	30.8	29.5	28.7	28.2	27.9	27.7	27.5	27.3	27.2	27.1	27.1	27.0	26.9	26.8
4	21.2	18.0	16.7	16.0	15.5	15.2	15.0	14.8	14.7	14.5	14.4	14.4	14.3	14.2	14.2
5	16.3	13.3	12.1	11.4	11.0	10.7	10.5	10.3	10.2	10.1	9.96	9.89	9.82	9.77	9.68
6	13.7	10.9	9.78	9.15	8.75	8.47	8.26	8.10	7.98	7.87	7.79	7.72	7.66	7.60	7.52
7	12.2	9.55	8.45	7.85	7.46	7.19	6.99	6.84	6.72	6.62	6.54	6.47	6.41	6.36	6.27
8	11.3	8.65	7.59	7.01	6.63	6.37	6.18	6.03	5.91	5.81	5.73	5.67	5.61	5.56	5.48
9	10.6	8.02	6.99	6.42	6.06	5.80	5.61	5.47	5.35	5.26	5.18	5.11	5.05	5.00	4.92
10	10.0	7.56	6.55	5.99	5.64	5.39	5.20	5.06	4.94	4.85	4.77	4.71	4.65	4.60	4.52
11	9.65	7.21	6.22	5.67	5.32	5.07	4.89	4.74	4.63	4.54	4.46	4.40	4.34	4.29	4.21
12	9.33	6.93	5.95	5.41	5.06	4.82	4.64	4.50	4.39	4.30	4.22	4.16	4.10	4.05	3.97
13	9.07	6.70	5.74	5.21	4.86	4.62	4.44	4.30	4.19	4.10	4.02	3.96	3.91	3.86	3.78
14	8.86	6.51	5.56	5.04	4.70	4.46	4.28	4.14	4.03	3.94	3.86	3.80	3.75	3.70	3.62
16	8.53	6.23	5.29	4.77	4.44	4.20	4.03	3.89	3.78	3.69	3.62	3.55	3.50	3.45	3.37
18	8.29	6.01	5.09	4.58	4.25	4.01	3.84	3.71	3.60	3.51	3.43	3.37	3.32	3.27	3.19
20	8.10	5.85	4.94	4.43	4.10	3.87	3.70	3.56	3.46	3.37	3.29	3.23	3.18	3.13	3.05
22	7.95	5.72	4.82	4.31	3.99	3.76	3.59	3.45	3.35	3.26	3.18	3.12	3.07	3.02	2.94
24	7.82	5.61	4.72	4.22	3.90	3.67	3.50	3.36	3.26	3.17	3.09	3.03	2.98	2.93	2.85
26	7.72	5.53	4.64	4.14	3.82	3.59	3.42	3.29	3.18	3.09	3.02	2.96	2.90	2.86	2.78
28	7.64	5.45	4.57	4.07	3.75	3.53	3.36	3.23	3.12	3.03	2.96	2.90	2.84	2.79	2.72
30	7.56	5.39	4.51	4.02	3.70	3.47	3.30	3.17	3.07	2.98	2.91	2.84	2.79	2.74	2.66
40	7.31	5.18	4.31	3.83	3.51	3.29	3.12	2.99	2.89	2.80	2.73	2.66	2.61	2.56	2.48
50	7.17	5.06	4.20	3.72	3.41	3.19	3.02	2.89	2.79	2.70	2.63	2.56	2.51	2.46	2.38
60	7.08	4.98	4.13	3.65	3.34	3.12	2.95	2.82	2.72	2.63	2.56	2.50	2.44	2.39	2.31
80	6.96	4.88	4.04	3.56	3.26	3.04	2.87	2.74	2.64	2.55	2.48	2.42	2.36	2.31	2.23
100	6.90	4.82	3.98	3.51	3.21	2.99	2.82	2.69	2.59	2.50	2.43	2.37	2.31	2.26	2.19
200	6.76	4.71	3.88	3.41	3.11	2.89	2.73	2.60	2.50	2.41	2.34	2.27	2.22	2.17	2.09
500	6.69	4.65	3.82	3.36	3.05	2.84	2.68	2.55	2.44	2.36	2.28	2.22	2.17	2.12	2.04
∞	6.63	4.61	3.78	3.32	3.02	2.80	2.64	2.51	2.41	2.32	2.25	2.18	2.13	2.08	2.00

* Multiply the numbers of the first row ($m_2 = 1$) by 10.

جدول آماری ۴- سطح زیر منحنی تابع چگالی F (ادامه)

$$F_{0.01, m_1, m_2} \text{ such that } P(F_{m_1, m_2} > F_{0.01, m_1, m_2}) = 0.01$$

$$= \int_{F_{0.01, m_1, m_2}}^{\infty} f(F) dF = 1 - \int_0^{F_{0.01, m_1, m_2}} f(F) dF$$

18	20	22	24	26	28	30	40	50	60	80	100	200	500	∞	m_1/m_2
619	621	622	623	624	625	626	629	630	631	633	633	635	636	637	1
99.4	99.4	99.5	99.5	99.5	99.5	99.5	99.5	99.5	99.5	99.5	99.5	99.5	99.5	99.5	2
26.8	26.7	26.6	26.6	26.6	26.5	26.5	26.4	26.4	26.3	26.3	26.2	26.2	26.1	26.1	3
14.1	14.0	14.0	13.9	13.9	13.9	13.8	13.7	13.7	13.7	13.6	13.6	13.5	13.5	13.5	4
9.61	9.55	9.51	9.47	9.43	9.40	9.38	9.29	9.24	9.20	9.16	9.13	9.08	9.04	9.02	5
7.45	7.40	7.35	7.31	7.28	7.25	7.23	7.14	7.09	7.06	7.01	6.99	6.93	6.90	6.88	6
6.21	6.16	6.11	6.07	6.04	6.02	5.99	5.91	5.86	5.82	5.78	5.75	5.70	5.67	5.65	7
5.41	5.36	5.32	5.28	5.25	5.22	5.20	5.12	5.07	5.03	4.99	4.96	4.91	4.88	4.85	8
4.86	4.81	4.77	4.73	4.70	4.67	4.65	4.57	4.52	4.48	4.44	4.42	4.36	4.33	4.31	9
4.46	4.41	4.36	4.33	4.30	4.27	4.25	4.17	4.12	4.08	4.04	4.01	3.96	3.93	3.91	10
4.15	4.10	4.06	4.02	3.99	3.96	3.94	3.86	3.81	3.78	3.73	3.71	3.66	3.62	3.60	11
3.91	3.86	3.82	3.78	3.75	3.72	3.70	3.62	3.57	3.54	3.49	3.47	3.41	3.38	3.36	12
3.72	3.66	3.62	3.59	3.56	3.53	3.51	3.43	3.38	3.34	3.30	3.27	3.22	3.19	3.16	13
3.56	3.51	3.46	3.43	3.40	3.37	3.35	3.27	3.22	3.18	3.14	3.11	3.06	3.03	3.00	14
3.31	3.26	3.22	3.18	3.15	3.12	3.10	3.02	2.97	2.93	2.89	2.86	2.81	2.78	2.75	16
3.13	3.08	3.03	3.00	2.97	2.94	2.92	2.84	2.78	2.75	2.70	2.68	2.62	2.59	2.57	18
2.99	2.94	2.90	2.86	2.83	2.80	2.78	2.69	2.64	2.61	2.56	2.54	2.48	2.44	2.42	20
2.88	2.83	2.78	2.75	2.72	2.69	2.67	2.58	2.53	2.50	2.45	2.42	2.36	2.33	2.31	22
2.79	2.74	2.70	2.66	2.63	2.60	2.58	2.49	2.44	2.40	2.36	2.33	2.27	2.24	2.21	24
2.72	2.66	2.62	2.58	2.55	2.53	2.50	2.42	2.36	2.33	2.28	2.25	2.19	2.16	2.13	26
2.65	2.60	2.56	2.52	2.49	2.46	2.44	2.35	2.30	2.26	2.22	2.19	2.13	2.09	2.06	28
2.60	2.55	2.51	2.47	2.44	2.41	2.39	2.30	2.25	2.21	2.16	2.13	2.07	2.03	2.01	30
2.42	2.37	2.33	2.29	2.26	2.23	2.20	2.11	2.06	2.02	1.97	1.94	1.87	1.83	1.80	40
2.32	2.27	2.22	2.18	2.15	2.12	2.10	2.01	1.95	1.91	1.86	1.82	1.76	1.71	1.68	50
2.25	2.20	2.15	2.12	2.08	2.05	2.03	1.94	1.88	1.84	1.78	1.75	1.68	1.63	1.60	60
2.17	2.12	2.07	2.03	2.00	1.97	1.94	1.85	1.79	1.75	1.69	1.66	1.58	1.53	1.49	80
2.12	2.07	2.02	1.98	1.94	1.92	1.89	1.80	1.73	1.69	1.63	1.60	1.52	1.47	1.43	100
2.02	1.97	1.93	1.89	1.85	1.82	1.79	1.69	1.63	1.58	1.52	1.48	1.39	1.33	1.28	200
1.97	1.92	1.87	1.83	1.79	1.76	1.74	1.63	1.56	1.52	1.45	1.41	1.31	1.23	1.16	500
1.93	1.88	1.83	1.79	1.76	1.72	1.70	1.59	1.52	1.47	1.40	1.36	1.25	1.15	1.00	∞

* Multiply the number of the first row ($m_2 = 1$) by 10.

Bendat J S & Piersol A G. Random Data: Analysis and Measurement Procedures. New York: Wiley, 1971. 407 p.

E.M. Mikhail & G. Gracie, Analysis and Adjustment of Survey Measurements, Wiley, John & Sons, Incorporated, October 1981, 368pp

Petr Vaníček, Edward J. Krakiwsky, Geodesy: the concepts, 2nd ed., North Holland, 1986, 697 p.

Moritz, H. Advanced least squares methods. Reports of the Department of Geodetic Science, No. 75, the Ohio State University, Columbus, 1972.

Goldman, S. Information Theory. New York: Dover, 1969, 397 p

E.M. Mikhail & F. Ackermann, 1976. Observations and Least Squares. Harper & Row, New York, 497p

Y. Djamour, 2004, Contribution de la Géodésie (GPS et Nivellement) a l'étude de la déformation tectonique et de l'alea sismique sur la région de Téhéran (Montagnes Alborz, Iran), université Montpellier II, Montpellier, 180pp, PhD thesis.

آریا نژاد، میربهادر قلی و ذهبیون، محمد، ۱۳۷۱، مقدمه‌ای بر آمار و احتمالات کاربردی، انتشارات دانشگاه علم و صنعت ایران، ۴۴۱ ص

پورپاک، علی محمد، ۱۳۸۷، محاسبات عددی: آنالیز عددی کاربردی، انتشارات جهاد دانشگاهی دانشگاه تهران دانشکده فنی، چاپ هفتم، ۴۰۰ ص

زالی، عباسعلی و جعفری شبستری، جمشی، ۱۳۸۵، مقدمه‌ای بر احتمالات و آمار، انتشارات دانشگاه تهران، چاپ هشتم، ۵۰۹ ص