

آمار و احتمال @toosarastat

فصل اول آمار توصیفی

مقدمه

امروزه در علوم پایه و مهندسی به دفعات نیاز به جمع آوری اطلاعات در مورد مجموعه‌های از اشیاء و یا انسانها داریم که این امر بر عهده علم آمار می‌باشد و با کمک آمار توصیفی می‌توان اطلاعات جمع آوری شده را به صورتی منظم گردآوری نمود بطوریکه بتوان با یک نگاه اجمالی به نتایج بدست آمده یک دید کلی نسبت به کل داده‌ها بدست آورد. در این فصل به چگونگی جمع آوری، اطلاعات و تجزیه و تحلیل آن با کمک نمودارها و پارامترهای مرکزی و پراکندگی می‌پردازیم.

۱-۱ تعاریف اولیه

جامعه آماری: به مجموعه‌ای از اشیاء یا افراد که حداقل یک ویژگی مشترک آنها مورد مطالعه قرار می‌گیرد. جامعه آماری می‌گوییم. ویژگی‌های مشترک یک جامعه آماری از عضوی به عضو دیگر تغییر می‌کند. که آنها را متغیر می‌نامند. متغیرها معمولاً به دو نوع تقسیم می‌شوند: متغیر کمی: متغیرهایی می‌باشند که معمولاً قابل اندازه‌گیری هستند و می‌توان مقدار آنها را به صورت عددی نمایش داد مثل مقدار وزن، قد، حجم و

۱- متغیر کیفی: متغیرهایی می‌باشند که مستقیماً توسط اعداد و ارقام قابل اندازه‌گیری نیستند. مثل گروه خونی، شغل، رنگ چشم و ... که برای اندازه‌گیری این متغیرها به آنها عددی نسبت می‌دهیم. متغیرهای کمی خود بر دو نوع هستند.

۱- گسسته: متغیرهایی که بین دو مقدار متصور آنها هیچ عدد دیگری وجود نداشته باشد.

۲- پیوسته: متغیرهایی که بین هر دو مقدار متصور آنها همواره عددی دیگری وجود دارد. مثل وزن یا طول و قد افراد.

پس از جمع آوری داده‌ها برای رسیدن به اهداف مورد نیاز به بررسی و تجزیه و تحلیل داده‌ها داریم که برای این منظور ابتدا داده‌ها را در یک جدول تنظیم و طبقه‌بندی می‌کنیم و سپس با استفاده از نمودارهای آماری نحوه توزیع داده‌ها را نمایش می‌دهیم و در نهایت داده‌ها را با کمک چند عدد به نام شاخص یا آماره خلاصه می‌کنیم.

۱-۲ جدول آماری

در جداول آماری علاوه بر داده‌ها، تعداد و درصد تکرار آنها نمایش داده می‌شود بطوریکه با یک نگاه به جدول می‌توان اطلاعات مفیدی در مورد پراکندگی و توزیع داده‌ها بدست آورد. یکی از متداول‌ترین جداول آماری جدول فراوانی است در جداول فراوانی همواره موارد زیر را خواهیم داشت: ۱- فراوانی: با شمردن تعداد دفعات تکرار هر داده در میان تمامی داده‌ها فراوانی آن داده بدست می‌آید. فرض کنید N داده داشته باشیم با قرار دادن داده‌های مشابه در یک دسته، در نهایت K طبقه خواهیم داشت ($K \leq N$) که تعداد داده‌ها در هر دسته را فراوانی آن داده می‌نامیم. فراوانی طبقه i

ام را با f_i نمایش می‌دهیم که $1 \leq i \leq k$ و $1 \leq f_i \leq N$ و $\sum_{i=1}^k f_i = N$ (تعداد کل داده‌ها)

۲- فراوانی نسبی: عبارتست از حاصل تقسیم فراوانی هر طبقه بر تعداد کل داده‌ها که آنرا با $r_i = \frac{f_i}{N}$ ، $i = 1, \dots, k$ نمایش می‌دهیم.

$$\sum_{i=1}^k r_i = \sum_{i=1}^k \frac{f_i}{N} = \frac{1}{N} \sum_{i=1}^k f_i = 1$$

هم چنین اگر فراوانی نسبی را در عدد ۱۰۰ ضرب کنید درصد وقوع هر داده را در میان کل داده‌ها بدست می‌آورید.

۳- فراوانی تجمعی نسبی: حاصل جمع فراوانی هر طبقه با طبقات قبل از آنرا فراوانی تجمعی می‌نامیم و به g_i ($1 \leq i \leq k$) نمایش می‌دهیم.

$$\text{فراوانی تجمعی طبقه } i = g_i = f_1 + f_2 + \dots + f_i = \sum_{j=1}^i f_j$$

به همین ترتیب حاصل جمع فراوانی نسبی هر طبقه با طبقات قبل آنرا فراوانی تجمعی نسبی می‌نامیم و به $(1 \leq i \leq k)$ s_i نمایش می‌دهیم

$$\text{فراوانی تجمعی طبقه } i = s_i = r_1 + r_2 + \dots + r_i = \sum_{j=1}^i s_j = s_i = \frac{1}{N} g_i$$

توجه: در محاسبه g_i و s_i اگر x_i نماینده طبقه i ام باشد می‌بایستی داده‌ها به صورت مرتب و از کوچک به بزرگ در جدول قرار داده شوند بطوریکه داشته باشیم $x_1 < x_2 < \dots < x_n$.

۱-۲-۱ جدول فراوانی برای داده های گسسته

در مثال زیر چگونگی تشکیل جدول فراوانی برای داده‌های گسسته را بیان می‌کنیم.

مثال ۱: داده‌های زیر تعداد فرزندان تحت پوشش بیمه در ۳۰ خانواده را نشان می‌دهد.

5, 4, 2, 1, 1, 1, 3, 3, 3, 2, 0, 5, 0, 2, 2, 2, 2, 4, 4, 1, 0, 1, 1, 3, 2, 5, 2, 1, 0, 3

ابتدا داده‌ها را از کوچک به بزرگ برای تشکیل جداول فراوانی مرتب می‌کنیم:

0, 0, 0, 0, 1, 1, 1, 1, 1, 1, 1, 2, 2, 2, 2, 2, 2, 2, 2, 3, 3, 3, 3, 4, 4, 4, 4, 5, 5, 5

در این صورت ۶ طبقه بدست خواهد آمد. نماینده هر طبقه در ستون x_i و فراوانی آن در ستون f_i نوشته می‌شود به این ترتیب جدول زیر را خواهیم داشت:

x_i	f_i	r_i	g_i	s_i
۰	۴	۰/۱۳	۴	۰/۱۳
۱	۷	۰/۲۳	۱۱	۰/۳۶
۲	۸	۰/۲۶	۱۹	۰/۶۲
۳	۴	۰/۱۳	۲۳	۰/۷۵
۴	۴	۰/۱۳	۲۷	۰/۸۸
۵	۳	۰/۱	۳۰	۱
جمع	۳۰	۱		

با توجه به جدول می‌توان به نتایج زیر رسید:

۱- با توجه به عدد ۰/۲۶ در ستون فراوانی نسبی می‌توان نتیجه گرفت که ۲۶ درصد از خانوارها دارای ۲ فرزند می‌باشند.

۲- با توجه به عدد ۰/۷۵ در ستون فراوانی تجمعی نسبی می‌توان نتیجه گرفت که ۷۵ درصد خانوارها حداکثر دارای ۳ فرزند می‌باشند.

۱-۲-۲ جدول فراوانی برای داده های پیوسته

مثال ۲: داده‌های زیر طول عمر ۴۰ عدد لامپ را نشان می‌دهند که به نزدیک‌ترین عدد صحیح گرد شده‌اند.

11 9 12 15 20 13 14 17 23 22
8 16 17 21 11 18 21 12 11 10
14 13 19 16 15 17 20 8 7 13
15 17 16 14 22 12 11 9 18 19

زمانی که با داده‌های پیوسته کار می‌کنیم داده‌ها را به رده‌های (فاصله‌ها) با طول مساوی تقسیم می‌کنیم و فراوانی داده‌ها را در هر رده بدست می‌آوریم. در این حالت نیاز به ثبت تمامی داده‌ها در جدول نمی‌باشد بلکه هر رده را به صورت مرتب از کوچک به بزرگ در جدول ثبت می‌کنیم. در این حالت روند به این صورت است که:

۱- از آنجا که هر عدد به نزدیکترین عدد صحیح گرد شده است پس مثلاً عدد ۱۲ در واقع عددی بین (۱۲/۵ و ۱۱/۵) بوده است برای اینکه این مقدار نیز در عملیات وارد شود از مفهوم میزان تغییر پذیری داده‌ها که با S نمایش می‌دهیم استفاده می‌کنیم.

$$S = \frac{\text{واحد گرد شده داده‌ها}}{2} = \frac{1}{2} = 0.5$$

۲- دامنه واقعی داده‌ها و کوچکترین و بزرگترین داده را به این ترتیب محاسبه می‌کنیم.

$$\min = \text{کوچکترین داده} - S = 7 - 0.5 = 6.5$$

$$\max = \text{بزرگترین داده} + S = 23 + 0.5 = 23.5$$

$$R = \max - \min = 23.5 - 6.5 = 17$$

۳- با توجه به دامنه داده‌ها می‌توان طول هر رده را معین نمود. برای این می‌بایستی دامنه را بر تعداد رده‌ها (که به صورت دلخواه قابل انتخاب است) تقسیم نمود.

$$\text{تعداد رده‌ها} = K \quad \omega = \frac{R}{K} = \text{طول رده}$$

معمولاً تعداد رده‌ها طوری انتخاب می‌شوند که هر رده حداقل پنج داده را در برداشته باشد می‌توان از فرمول $K = 1 + 3.322 \log_{10}^N$ برای تعیین تعداد رده استفاده نمود بنابر این:

$$K = 1 + 3.322 \log_{10}^{40} = 6.322 \cong 7 \Rightarrow \omega = \frac{17}{7} = 2.4 \cong 3$$

۴- حالا کافیست ۷ رده به طول ۳ در ستون اول جدول فراوانی ایجاد کنیم و مجدداً مطابق مثال قبل جدول را کامل کنیم. اولین رده از کوچکترین عنصر آغاز می‌شود.

رده‌ها	خط و نشان	x_i	f_i	r_i	g_i	s_i
6/5-9/5		۸	۵	0/125	۵	0/125
9/5-12/5		۱۱	۸	0/2	۱۳	0/325
12/5-15/5		۱۴	۹	0/225	۲۲	0/55
15/5-18/5		۱۷	۹	0/225	۳۱	0/775
18/5-21/5		۲۰	۶	0/15	۳۷	0/925
21/5-24/5		۲۳	۳	0/075	۴۰	1/0
جمع			۴۰			

توجه شماره ۱ :

- X_i معرف نماینده هر رده است که عضو میانی رده می‌باشد.

- چون آخرین رده شامل هیچ عضوی نبود آنرا حذف می‌کنیم مثلاً در این مثال رده ۲۴/۵-۲۷/۵ شامل هیچ عضوی نیست بنابر این حذف می‌شود.

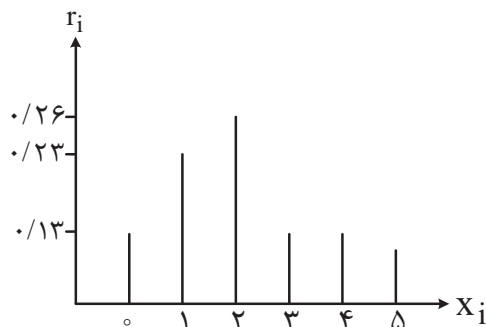
۳-۱ نمودارهای آماری

یکی دیگر از روشهای مناسب برای خلاصه نمودن داده‌های آماری استفاده از نمودار می‌باشد. نمودارها بهترین ابزار برای نمایش نحوه توزیع داده‌ها می‌باشند که در ذیل معروفترین آنها را معرفی می‌کنیم.

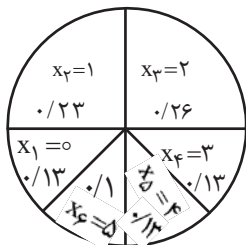
۱-۳-۱ نمودارهای آماری برای داده‌های گسسته

۱- نمودارهای میله‌ای:

در این نمودار محور X ها نمایش دهنده مقادیر داده‌ها و محور Y ها نمایش دهنده فراوانی نسبی می‌باشد. چون فراوانی نسبی داده‌های هر جامعه آماری مقادیر بین صفر و یک را می‌گیرد بنابر این می‌توان چندین نمودار را با یکدیگر مقایسه نمود. شکل نمودار میله‌ای را برای مثال ۱ نمایش می‌دهد.



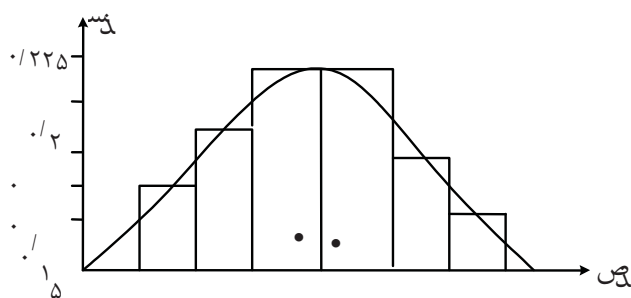
۳- نمودار دایره‌ای: این نمودار درون دایره‌ای رسم می‌شود که به قطاع‌هایی تقسیم شده است و هر قطاع متناسب با مساحت خود معرف یکی از فراوانی‌های نسبی در جدول فراوانی می‌باشد. شکل زیر نمایش دهنده نمودار دایره‌ای برای مثال ۱ می‌باشد.



۱-۳-۲ نمودارهای آماری برای داده‌های پیوسته

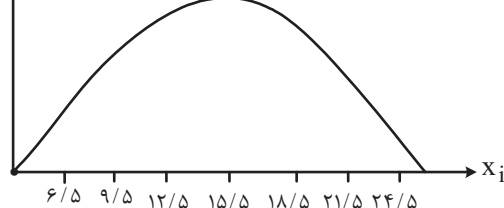
همانطور که در جدول فراوانی مشاهده کردید تفاوت عمده رده‌های پیوسته با گسسته در بکارگیری رده‌ها می‌باشد این امر در نمودارها نیز کاملاً صادق است. در این قسمت سه نمودار را که برای نمایش داده‌های پیوسته بکار می‌روند معرفی می‌کنیم.

۱- هیستوگرام (نمودار ستونی): این نمودار کاملاً مشابه نمودار میله‌ای می‌باشد با این تفاوت که محور X ها نمایش دهنده رده‌های جدول فراوانی است و به ازای هر رده مستطیلی که عرض آن یک واحد و ارتفاع آن معادل فراوانی نسبی آن رده است رسم می‌شود. بنابر این مجموع مساحت‌های مستطیل‌های رسم شده برابر واحد می‌باشد که همان مجموع کل فراوانی نسبی‌ها است. شکل زیر نمودار هیستوگرام برای مثال ۲ می‌باشد.



۲- چند برابر فراوانی: اگر در نمودار هیستوگرام وسط قاعده‌های بالای مستطیل‌ها را به یکدیگر توسط خطوطی به صورت متوالی متصل کنیم نمودار چند برابر فراوانی بدست می‌آید که در شکل بالا نیز نشان داده شده است. ابتدای خطوط به وسط رده ماقبل و انتهای خطوط به وسط رده مابعد مستطیل‌های هیستوگرام متصل می‌شوند به این ترتیب مساحت زیر نمودار چند بر فراوانی مجدداً برابر واحد خواهد بود.

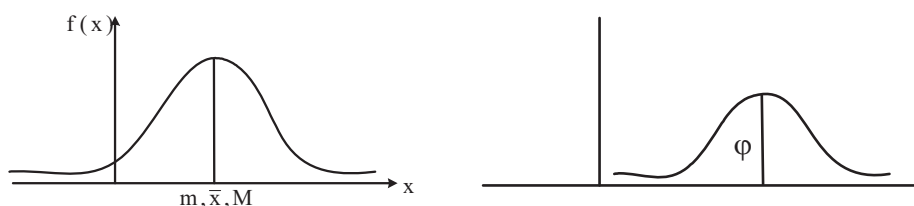
۳- منحنی فراوانی: در صورتی که تعداد داده‌ها زیاد باشد و طول رده‌ها کوچک، تعداد رده‌ها زیاد می‌شود و در نتیجه تعداد اضلاع چند بزرگ‌افراوانی افزایش یافته و در نهایت مشابه یک منحنی خواهد شد که در این حالت نیز مساحت زیر منحنی برابر واحد است. شکل زیر نمونه‌ای از فراوانی است.



متغیر تصادفی نرمال: اگر مجموعه داده‌های X دارای میانگین $\bar{X}$ و واریانس σ^2 باشند و نمودار آنها از تابع

$$f(x) = \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(x-\bar{x})^2}{2\sigma^2}} \quad (-\infty < x < +\infty)$$

تبعیت کند می‌گوییم متغیر تصادفی X نرمال می‌باشد. نمودار متغیر تصادفی نرمال نسبت به خط $x = \bar{X}$ متقارن می‌باشد. به شکل زیر توجه کنید.



همچنین با افزایش مقدار δ نمودار در راستای محور x ها کشیده تر می‌شود.

با افزایش مقدار $\bar{X}$ نمودار به سمت راست و با کاهش آن به سمت چپ جابجا می‌شود.

۴-۱ پارامترهای مرکزی و پراکندگی

از آنجا که با مطالعه یک جامعه آماری تعداد زیادی داده بدست می‌آوریم و مطالعه روی تک تک یا قسمتی از این داده‌ها مشکل و حتی غیر ممکن است همواره علاقه داریم این داده‌ها را با کمک شاخص‌ها و پارامترهایی خلاصه کنیم تا بتوان با یک نگاه اجمالی به آن یک دید کلی نسبت به کل داده‌ها و جامعه آماری بدست آورد. بنابر این پارامترهای مرکزی و پراکندگی را معرفی می‌کنیم.

۱-۴-۱ پارامترهای مرکزی

عموماً داده‌ها در جامعه آماری یکنوع تجمع و فشردگی حول یک مقدار خاص از صفت مورد مطالعه را بوجود می‌آورند که این مقدار خاص به عنوان یک پارامتر مرکزی معرفی می‌شود. مهمترین پارامترهای مرکزی عبارتند از: میانگین، میانه و مد (نما)

۱- میانگین: میانگین بر چند نوع است که معروفترین آنها میانگین حسابی، هندسی، همساز (هارمونیک) و درجه دوم می‌باشد.

الف) میانگین حسابی: اگر داده‌های $x_1, x_2, \dots, x_k$ به ترتیب دارای فراوانی $f_1, f_2, \dots, f_k$ باشند و تعداد کل این داده‌ها N باشد $\bar{X}$ یا میانگین حسابی به صورت زیر تعریف می‌شود:

$$\bar{X} = \frac{1}{N} \sum_{i=1}^K f_i x_i$$

برای داده‌های پیوسته X_i را نماینده هر رده در نظر می‌گیریم. به عنوان مثال میانگین حسابی برای مثال ۱ و ۲ عبارتست از:

$$\text{مثال 1: } \bar{X} = \frac{1}{30} (4 \times 0 + 7 \times 1 + 8 \times 2 + 4 \times 3 + 4 \times 4 + 3 \times 5) = 2/2$$

$$\text{مثال 2: } \bar{X} = \frac{1}{40} (8 \times 5 + 11 \times 8 + 14 \times 9 + 17 \times 9 + 20 \times 6 + 23 \times 3) = 14/9$$

ب) میانگین هندسی: در صورتی که همگی داده‌ها مثبت باشند میانگین هندسی به صورت زیر محاسبه می‌شود.

$$G = \left(\prod_{i=1}^N x_i^{f_i} \right)^{\frac{1}{N}}$$

ج: میانگین همساز: (همنوا یا هارمونیک): اگر هیچکدام از داده ها صفر نباشد می توان میانگین همساز را از طریق فرمول محاسبه نمود.

$$H = \frac{N}{\sum_{i=1}^n \frac{f_i}{x_i}}$$

در حالت کلی رابطه $H \leq G \leq \bar{X}$ بین این سه میانگین برقرار است. که حالت تساوی زمانی رخ می دهد که همگی داده ها برابر باشند.

۲- میانه: با مرتب نمودن داده ها به صورت غیر نزولی عدد m را بعنوان میانه آنها در نظر می گیریم اگر تقریباً نیمی از داده ها سمت چپ آن و نیمی دیگر در سمت راست آن قرار داشته باشند.

همچنین با توجه به جدول فراوانی، میانه کوچکترین مقدار x است که فراوانی تجمعی آن بیشتر یا مساوی با $\frac{N}{2}$ باشد.

- برای داده های گسسته در صورتی که داده ها را به صورت غیر نزولی مرتب کنیم اگر تعداد داده ها فرد باشد میانه داده وسطی یا به عبارتی $m = x_{(\frac{n+1}{2})}$ خواهد بود و اگر تعداد داده ها زوج باشد، میانه میانگین دو داده وسطی می باشد

$$m = \frac{x_{\left(\frac{n}{2}\right)} + x_{\left(\frac{n}{2}+1\right)}}{2}; m = \frac{x_{\left(\frac{n}{2}\right)} + x_{\left(\frac{n}{2}+1\right)}}{2} = \frac{2+2}{2} = 4$$

- برای داده های پیوسته ابتدا می بایستی رده میانه دار را پیدا نمود. رده میانه دار اولین رده ای است که فراوانی تجمعی نسبی آن از 0.5 بیشتر است. میانه عددیست که در این رده قرار دارد. حال برای محاسبه آن از فرمول زیر استفاده می کنیم:

$$m = L_{0/5} + \frac{(0.5N - g_{0/5}) \omega}{f_{0/5}}$$

که در آن :

$L_{0/5}$: کران پایین رده میانه دار

N : تعداد داده ها

$g_{0/5}$: فراوانی تجمعی رده قبل از رده میانه دار

$f_{0/5}$: فراوانی رده میانه دار

ω : طول رده

به عنوان مثال میانه برای مثال ۱ با توجه به تعداد زوج داده ها عبارتست از:

برای مثال ۲ نیز میانه برابر است با:

$$m = 12/5 + \frac{(0.5 \times 40 - 13)3}{9} = 14/8$$

۳- مد یا نما: داده ای که فراوانی آن از سایر داده ها بیشتر باشد مد یا نما گفته می شود و با M نمایش داده می شود.

الف) روش محاسبه مد برای داده های گسسته: پس از بدست آوردن فراوانی داده، داده های که فراوانی آن از بقیه بیشتر باشد مد می باشد در این صورت سه حالت زیر پیش می آید: (که در مثال زیر به آن می پردازیم)

مثال ۳: برای داده های مثال ۱ تنها عدد ۲ دارای بیشترین فراوانی می باشد. پس عدد ۲ مقدار مد خواهد بود. اما برای داده های ۵ و ۴ و ۴ و ۲ و ۲ و ۱ و ۱ و ۱ چون فراوانی ۱ و ۲ برابر است و این دو عدد به صورت متوالی قرار گرفته اند در این حالت مد برابر است با میانگین این دو عدد:

$$M = \frac{1+2}{2} = \frac{3}{2}$$

اما برای داده‌های ۴ و ۲ و ۱ و ۱، فراوانی ۴ مساوی می‌باشد ولی مجاور یکدیگر نیستند دو مقدار برای مد خواهیم داشت که عبارتند از:

$$M_1 = 1 \quad M_2 = 4$$

توجه کنید در صورتی که فراوانی همه داده‌ها برابر باشد داده‌ها بدون مد در نظر گرفته می‌شوند.

ب) محاسبه مد برای داده‌های پیوسته: رده‌ای که فراوانی آن از سایر رده‌ها بیشتر است را به عنوان رده نمایی انتخاب می‌کنیم در این حالت می‌توان نماینده رده را به عنوان مد یا نما انتخاب کرد. اما برای محاسبه دقیق تر از فرمول زیر نیز می‌توان استفاده نمود:

$$M = L_M + \left(\frac{D_1}{D_1 + D_2} \right) \omega$$

L_M : کران پایین رده نمایی.

D_1 : اختلاف فراوانی‌های نسبی رده نمایی و رده قبل از آن.

D_2 : اختلاف فراوانی‌های نسبی رده نمایی و رده بعد از آن.

ω : طول رده.

در مثال ۲ داریم:

$$M_1 = 14 \quad ; \quad M_2 = 17$$

$$M = \frac{14+17}{2} = 15.5$$

و یا بصورت دقیق تر:

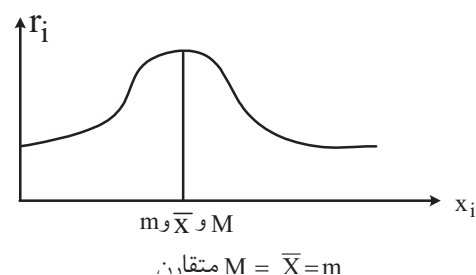
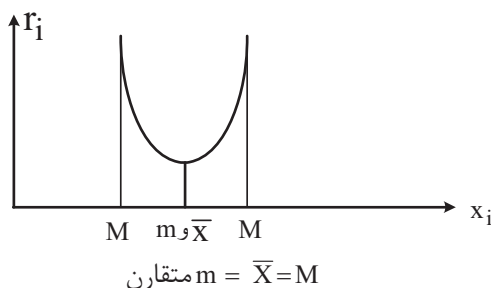
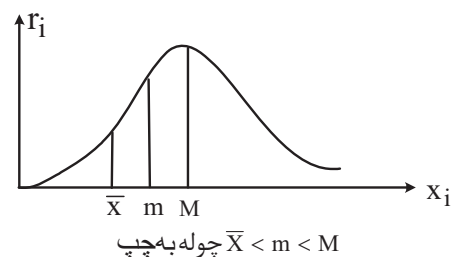
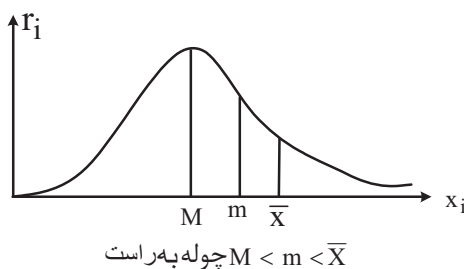
$$M_1 = 12.5 + \left(\frac{0.025}{0.025+0} \right) \times 3 = 15.5$$

$$M_2 = 15.5 + \left(\frac{0}{0+0.075} \right) \times 3 = 15.5$$

$$\Rightarrow M = 15.5$$

۱-۵ چولگی توزیع فراوانی

با رسم نمودار فراوانی پیوسته عموماً یکی از اشکال زیر بدست می‌آید. که چگونگی قرار گرفتن میانگین و میانه و مد را در اشکال زیر مشاهده می‌شود.



برای محاسبه میزان عدم تقارن منحنی، شاخصی به نام B_1 (ضریب چولگی) وجود دارد که در فصل‌های بعدی با آن آشنا می‌شوید. در توزیع‌هایی که چولگی زیاد نباشد رابطه تجربی زیر که به رابطه پیرسن معروف است برقرار می‌باشد.

$$\bar{X} - M \cong 3(\bar{X} - m)$$

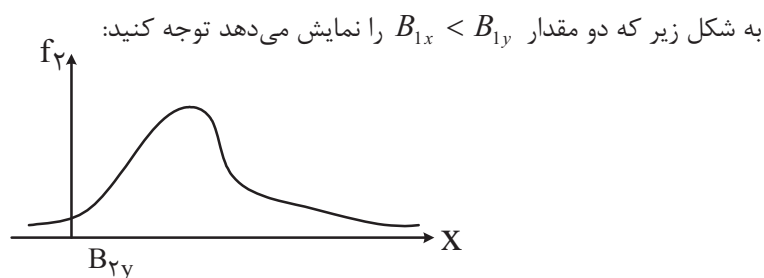
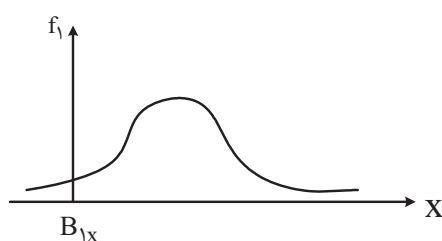
مقدار عدم تقارن منحنی که با ضریب چولگی سنجیده می‌شود را با نمودار نرمال مقایسه می‌کنیم برای بدست آوردن ضریب چولگی از مفهوم گشتاور مرتبه k ام طول میانگین (گشتاور مرکزی مرتبه k ام) استفاده می‌کنیم. که عبارتست از :

$$\mu_k = \frac{1}{N} \sum_{i=1}^n f_i (x_i - \bar{x})^k$$

ضریب چولگی که با B_1 نمایش داده می‌شود میزان عدم تقارن منحنی را نسبت به منحنی نرمال نمایش می‌دهد که به شکل زیر تعریف می‌شود:

$$B_1 = \frac{\mu_3}{\sqrt{\mu_2^3}} = \frac{\mu_3}{\delta^3}$$

در صورتی که منحنی متقارن باشد مقدار B_1 برابر صفر می‌باشد و هر چه B_1 مقدار بزرگتری داشته باشد نشان دهنده افزایش عدم تقارن می‌باشد.

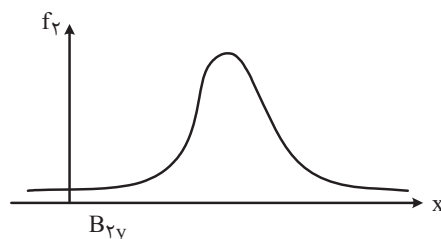
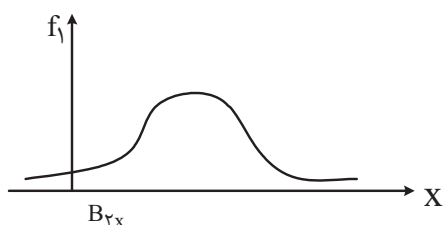


ضریب کشیدگی : نشان دهنده میزان کشیدگی منحنی می‌باشد و به صورت زیر تعریف می‌شود.

$$B_2 = \frac{\mu_4}{\mu_2^2} = \frac{\mu_4}{\delta^4}$$

دو منحنی زیر متقارن می‌باشند اما ضریب کشیدگی در آنها متفاوت است.

$$B_{2x} < B_{2y}$$



۱-۶ پارامترهای پراکندگی

شاخص های مرکزی تنها یک منطقه را به عنوان محل تمرکز داده ها معرفی می کنند حال آنکه ممکن است دو دسته داده با پراکندگی های متفاوت دارای میانگین برابری باشند به عبارتی نیاز به شاخصی برای نمایش میزان پراکندگی داده ها خواهیم داشت که در ادامه به معرفی این پارامترها می پردازیم.

۱- دامنه داده ها: حاصل تفاضل بزرگترین داده از کوچکترین داده را R یا دامنه داده ها گویند.

$$R = X_N - X_1$$

که در آن داده ها به فرم $X_{(1)} < X_{(2)} < \dots < X_N$ با فرض اینکه مرتب شده اند.

این شاخص معیار خوبی برای محاسبه میزان پراکندگی داده ها نیست. زیرا در محاسبه تنها کوچکترین و بزرگترین داده وارد می شود.

مثال ۴: نمرات ۱۰ دانش آموز در دو کلاس متفاوت در درس ریاضی به قرار زیر است:

کلاس A: ۰, 4, 4, 12, 12, 12, 14, 16, 16, 20

کلاس B: 8, 8, 9, 9, 12, 12, 12, 13, 13, 14

با محاسبه مقادیر $\bar{X}$ و m و M دیده می شود که برای هر دو کلاس داریم:

$$\bar{X} = 11, \quad M = 12, \quad m = 12$$

اما با توجه به مقادیر تک تک نمرات واضح است که میزان پراکندگی نمرات در دو کلاس کاملاً متفاوت است و برای این منظور نیاز به شاخص های پراکندگی می باشد تا این مطلب را بتوان با مقایسه آنها نشان داد.

برای دو کلاس مقدار دامنه را محاسبه می کنیم:

$$\text{کلاس A: } R = 20 - 0 = 20$$

$$\text{کلاس B: } R = 14 - 8 = 6$$

۲- میانگین انحرافات: فاصله داده X_i از میانگین را انحراف از میانگین داده X_i گویند که به صورت $|X_i - \bar{X}|$ محاسبه می شود. اگر این مقدار را

برای تمامی داده ها محاسبه کنیم و از نتیجه میانگین بگیریم میانگین انحرافات بدست خواهد آمد که عبارتست از: $D = \frac{1}{N} \sum_{i=1}^K f_i |X_i - \bar{X}|$

از آنجا که میانگین انحرافات به تمام داده ها وابسته است معیار مناسبی برای سنجش پراکندگی داده ها محسوب می شود اما بدلیل وجود قدر مطلق در فرمول، محاسبه آن مشکل است و نمی توان آنرا ساده نمود بنابر این از واریانس و انحراف استاندارد استفاده می کنیم.

۳- واریانس و انحراف استاندارد: میانگین مجذور انحرافات را واریانس می نامیم و با نماد σ^2 یا S_b^2 نمایش می دهیم. که عبارتست از:

$$S_b^2 = \frac{1}{N} \sum_{i=1}^K f_i (X_i - \bar{X})^2$$

در مبحث استنباط آماری واریانس را از مجموع مجذور انحرافات داده ها تقسیم بر $N-1$ بدست می آورند و آنرا با S^2 نمایش می دهند.

$$S^2 = \frac{1}{N-1} \sum_{i=1}^K f_i (X_i - \bar{X})^2$$

در مباحث این درس هر جا صحبت از واریانس می کنیم منظور S^2 می باشد.

اگر از واریانس جذر بگیریم یعنی $S = \sqrt{S^2}$ در این صورت S را انحراف استاندارد می نامیم که معیار مناسبی برای سنجش پراکندگی می باشد.

همچنین S^2 را می توان از فرمول زیر محاسبه نمود:

$$S^2 = \frac{1}{N-1} \left[\sum_{i=1}^K f_i X_i^2 - \frac{1}{n} \left(\sum_{i=1}^K f_i X_i \right)^2 \right] = \frac{1}{n-1} \left[\sum_{i=1}^K f_i X_i^2 - n \bar{X}^2 \right]$$

اثبات:

$$S^2 = \frac{1}{N-1} \sum_{i=1}^N f_i (x_i - \bar{x})^2 = \frac{1}{N-1} \sum_{i=1}^n f_i (x_i^2 - 2\bar{x}_i \bar{x} + \bar{x}^2)$$

$$= \frac{1}{N-1} \left(\sum_{i=1}^n f_i x_i^2 \right) - 2n\bar{x}^2 + n\bar{x}^2 = \frac{1}{N-1} \left(\sum_{i=1}^n f_i x_i^2 \right) - n\bar{x}^2$$

مثال ۵: مقدار واریانس را برای مثال ۲ محاسبه کنید:

$$\bar{X} = 14/9 \quad N=40$$

$$S^2 = \frac{1}{40-1} \left[(50 \times 64 + 8 \times 121 + 9 \times 196 + 9 \times (17)^2 + 6 \times 40 + 3 \times (23)^2 - 40 \times (14/9)^2) \right]$$

۴- ضریب تغییرات: در محاسبه میزان پراکندگی داده‌ها همواره با داده‌هایی سروکار داریم که با مقیاس‌های مختلفی اندازه‌گیری شده‌اند بنابراین برای مقایسه میزان پراکندگی داده‌های بدست آمده از دو جامعه آماری که با مقیاس‌های مختلفی اندازه‌گیری شده‌اند استفاده از واریانس مناسب نمی‌باشد

زیرا واریانس به مقیاس اندازه‌گیری وابسته می‌باشد بنابراین این از مقیاس مناسبتری به نام ضریب تغییرات استفاده می‌کنیم که از رابطه $CV = \frac{S}{\bar{X}}$ بدست می‌آید و معمولاً با ضریب آن در عدد صد بر حسب درصد بیان می‌شود.

مثال ۶: یک کارخانه تولید لاستیک دو نوع محصول A و B تولید می‌کند. لاستیک نوع A دارای میانگین طول عمر ۲۰۰۰ کیلو متر و انحراف استاندارد ۲۰۰ کیلو متر می‌باشد و لاستیک نوع B دارای میانگین طول عمر ۱۸۰۰ کیلو متر و انحراف استاندارد ۲۰۰ کیلو متر می‌باشد، کدام نوع لاستیک برای خرید مناسب‌تر می‌باشد؟

$$X_A = 2000 \quad \bar{X}_B = 1800 \quad \Rightarrow \quad CV_A = \frac{200}{2000} = 0.1$$

$$S_A = 200 \quad S_B = 200 \quad \Rightarrow \quad CV_B = \frac{200}{18000} = 0.01$$

$$CV_A = 0.1 \times 100 = 10\%$$

$$CV_B = 0.01 \times 100 = 1\%$$

همانطور که ملاحظه می‌کنید میانگین طول عمر لاستیک دوم از لاستیک اول کمتر است ولی با توجه به اینکه ضریب تغییرات لاستیک دوم کمتر از لاستیک اول است خرید لاستیک دوم به صرفه‌تر می‌باشد.

۷-۱ تغییر مقیاس و مبدأ

داده‌ها را با واحدهای متفاوتی می‌توان از جامعه آماری جمع آوری نمود به عنوان مثال فرض کنید داده‌های مربوط به وزن ۴۰ نفر از دانشجویان یک کلاس را با واحد کیلوگرم جمع آوری کرده باشید و بخواهید مقادیر میانگین و واریانس را بر حسب پاوند بدست بیاورید برای این منظور نیازی به محاسبه مجدد میانگین و واریانس نمی‌باشد بلکه کفایت از روش تغییر مقیاس و مبدأ استفاده کنید.

۱-۷ تغییر مقیاس

اگر تمامی داده‌ها در عدد a ضرب شوند در این صورت داریم:

$$x_1, x_2, \dots, x_n \rightarrow ax_1, ax_2, \dots, ax_n$$

$$\Rightarrow \bar{X} = \frac{1}{N} \sum_{i=1}^n x_i \rightarrow \bar{X}_a = \frac{1}{N} \sum_{i=1}^n ax_i = a \left(\frac{1}{N} \sum_{i=1}^n x_i \right) = a\bar{X} \Rightarrow \bar{X}_a = a\bar{X}$$

به همین ترتیب برای محاسبه واریانس بدست می‌آید:

تمرین ۱:

$$S_a^2 = a^2 S^2 \rightarrow S_a = |a| S$$

$$S_a^2 = \frac{1}{n-1} \sum_{i=1}^n (a x_i - a \bar{x})^2 = a^2 \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{x})^2 = a^2 S^2$$

$$\Rightarrow S_a = \sqrt{a^2 S^2} = |a| S$$

تمرین ۲:

$$X \rightarrow X + b$$

$$\bar{X}_b = \bar{X} + b$$

$$\bar{X}_b = \frac{1}{n} \sum_{i=1}^n (x_i + b) = \frac{1}{n} \left[\sum_{i=1}^n x_i + \sum_{i=1}^n b \right] = \frac{1}{n} \sum_{i=1}^n x_i + \frac{b}{n} \sum_{i=1}^n (1) = \bar{X} + b \frac{n}{n} = \bar{X} + b$$

$$S_b^2 = S^2$$

$$S_b^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i + b - \bar{X} - b)^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2 = S^2$$

تمرین ۳ -

$$X \rightarrow Y = aX + b$$

$$\bar{Y} = a\bar{X} + b, \quad S_Y^2 = a^2 S^2$$

$$\text{اثبات: } \bar{Y} = \frac{1}{n} \sum_{i=1}^n (ax_i + b) = \frac{a}{n} \sum_{i=1}^n x_i + \frac{bn}{n} = a\bar{X} + b$$

$$S_Y^2 = \frac{1}{n-1} \sum_{i=1}^n (ax_i + b - a\bar{X} - b)^2 = \frac{a^2}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2 = a^2 S^2$$

استاندارد سازی

مثال: نمره علی از امتحان فیزیک و ریاضی به ترتیب برابر ۴۰ و ۶۰ شده است اگر میانگین نمرات امتحان فیزیک و ریاضی به ترتیب برابر ۲۰ و ۵۰ باشد و انحراف معیار امتحان فیزیک و ریاضی به ترتیب برابر ۱ و ۲ باشد علی کدام درس را بهتر امتحان داده است.
حل: برای اینکه بتوان نمرات دو درس را با یکدیگر مقایسه نمود می‌بایستی ابتدا نمرات را استاندارد سازی نمود و سپس آنها را با یکدیگر مقایسه نمود.

$$\text{نمرات استاندارد شده علی در درس ریاضی} = \frac{60 - 50}{2} = 5$$

$$\text{نمرات استاندارد شده علی در درس فیزیک} = \frac{40 - 20}{1} = 20$$

با وجود اینکه نمره علی در درس فیزیک کمتر از ریاضی می‌باشد اما با استاندارد نمودن نمره دو درس مشاهده می‌کنیم که نمره وی در درس فیزیک بالاتر از درس ریاضی می‌باشد به عبارتی علی درس فیزیک را بهتر از درس ریاضی امتحان داده است.

۱-۷-۲ تغییر مبدأ

در صورتی که به تمام داده‌ها مقدار b را اضافه یا کم کنیم می‌توان نشان داد که مقادیر $\bar{X}$ و S^2 جدید از روابط زیر محاسبه می‌شوند:

$$\bar{X}_b = \bar{X} + b$$

$$S_b^2 = S^2 \quad \text{تغییر مبدأ روی واریانس بی تأثیر است}$$

با اعمال همزمان تغییر مبدأ و مقیاس خواهیم داشت:

$$\bar{Y} = a \bar{X} + b$$

$$S_a^2 = a^2 S^2$$

مطالب فوق برای میانه و مد نیز صادق می‌باشند و داریم:

$$m^1 = am + b$$

$$M^1 = a m + b$$

۱۳-۱ استاندارد سازی

از یک جامعه آماری n نمونه $X_1, X_2, \dots, X_n$ بصورت تصادفی انتخاب می‌کنیم بطوریکه میانگین و واریانس نمونه‌ها بترتیب $\bar{X}$ و S_X^2 می‌باشد. با توجه به تغییر مبدأ و مقیاس مقدار هر نمونه را از میانگین نمونه‌ها کم می‌کنیم و حاصل را بر S_X تقسیم می‌کنیم بنابر این داده‌های

$$y_1 = \frac{X_1 - \bar{X}}{S_X}, \quad y_2 = \frac{X_2 - \bar{X}}{S_X}, \quad y_n = \frac{X_n - \bar{X}}{S_X} \quad \text{را خواهیم داشت.}$$

$$\bar{y} = \frac{\bar{X} - \bar{X}}{S_X} = 0; \quad S_y^2 = \left(\frac{1}{S_X^2}\right) S_X^2 = 1 \quad \text{با محاسبه مقایر میانگین و واریانس داده‌های جدید داریم:}$$

همانطور که ملاحظه می‌کنید داده‌های جدید دارای میانگین صفر و واریانس ۱ می‌باشند که به آنها داده‌های استاندارد شده می‌گوییم. همینطور اگر

داده‌های X_1 تا X_n را با متغیر تصادفی X نمایش دهیم در این صورت $Y = \frac{X - \bar{X}}{S_X}$ فرم استاندارد شده یا صورت معیاری متغیر تصادفی X می‌باشد.

مسائل فصل اول :

۱- دو جامعه با اندازه‌های میانگین $\bar{X}_1, \bar{X}_2$ و انحراف معیار S_1, S_2 را با یکدیگر ادغام می‌کنیم ثابت کنید میانگین و انحراف معیار جدید از روابط زیر بدست می‌آید:

$$\bar{X} = \frac{N_1 \bar{X}_1 + N_2 \bar{X}_2}{N_1 + N_2}$$

$$S^2 = \frac{N_1 S_1^2 + N_2 S_2^2}{N_1 + N_2} + \frac{N_1 N_2}{(N_1 + N_2)^2} (\bar{X}_1 - \bar{X}_2)^2$$

۲- میانگین و واریانس ۲ داده به ترتیب ۱۵ و ۵ می‌باشد. اگر به جای عدد ۲۵ اشتباهاً عدد ۱۵ را در محاسبات اعمال کرده باشیم میانگین و واریانس جدید را بدست بیاورید.

۳- نشان دهید تغییر مقیاس داده بر روی مقدار ضریب تغییرات $CV = \frac{S}{\bar{X}}$ بی‌اثر می‌باشد، آیا این مطلب در مورد تغییر مبدأ نیز صادق است؟

۴- ثابت کنید برای میانگین حسابی، هندسی و همساز رابطه زیر برقرار است.

$$\bar{X}_H \leq \bar{X}_G \leq \bar{X}$$

۵- اگر میانگین را از داده‌های یک جامعه آماری کم کنیم و نتیجه را بر انحراف معیار تقسیم کنیم (یعنی $y_i = \frac{x_i - \bar{X}}{S}$) نشان دهید میانگین و انحراف معیار جدید به ترتیب صفر و یک می‌باشد.

۶- برای بدست آوردن یک معیار پراکندگی جدید داده‌ها را دو به دو با یکدیگر مقایسه می‌کنیم و میانگین n^2 داده جدید $(X_i - X_j)$ را با S_{ij}^2 نمایش می‌دهیم:

$$S_{ij}^2 = \frac{1}{n^2} \sum_{i=1}^n \sum_{j=1}^n (X_i - X_j)^2$$

نشان دهید $S_{ij}^2 = 2S^2$ (راهنمایی: داخل پراکنش مقدار $\bar{X}$ را اضافه و کم کنید)

۷- نشان دهید میانگین حسابی و واریانس نخستین n عدد طبیعی به ترتیب $\frac{n+1}{2}$ و $\frac{n^2-1}{12}$ می‌باشد.

۸- جدول زیر را برای داده‌ها و فراوانی آنها در نظر بگیرید:

x	۰	۱	۲	...	n
f	$\binom{n}{0}$	$\binom{n}{1}$	$\binom{n}{2}$	...	$\binom{n}{n}$

نشان دهید میانگین و واریانس این داده‌ها به ترتیب عبارتست از:

$$\bar{X} = \frac{n}{2}, \quad S^2 = \frac{n}{4}$$

۹- اگر متحرکی مسافت X_1 را با سرعت V_1 و ... و مسافت X_n را با سرعت V_n طی کنید ثابت کنید سرعت متوسط این متحرک با استفاده از رابطه میانگین همساز یا هارمونیک بدست می‌آید که در آن V_i معادل مقادیر داده‌ها و X_i معادل فراوانی آنهاست.

۱۰- عدد Q_P که $0 < P < 1$ را چندک P م داده‌ها تعریف می‌کنیم هر گاه فراوانی تجمعی نسبی (I_i) آن بزرگتر یا مساوی با عدد P باشد. به عبارت دیگر هر گاه XP ۱۰٪ از داده‌ها قبل از آن قرار گیرند. به عنوان مثال $Q_{.5}$ که به آن چارک دوم می‌گوییم همان میانه می‌باشد چرا که ۵۰٪ داده‌ها قبل از آن قرار دارند. در حالت کلی $Q_{.25}$ و $Q_{.5}$ و $Q_{.75}$ را به ترتیب با Q_1, Q_2, Q_3 نمایش می‌دهیم و به آنها چارکهای اول و دوم و سوم می‌گوییم.

الف) اگر برای داده‌های گسسته x_i داشته باشیم $x_1 < x_2 < \dots < x_i < \dots < x_N$ آنگاه نشان دهید $Q_P = (1-\omega) x_r + \omega x_{r+1}$ که در آن $\omega = (n+1)P - r$, $r = (n+1)P$.

ب) برای داده‌های پیوسته رده‌ای که فراوانی تجمعی آن بزرگتر یا مساوی با عدد P باشد را رده Q_P می‌نامیم نشان دهید چندک P م برای داده‌های پیوسته از رابطه زیر بدست می‌آید:

$$Q_P = L_p + \frac{(np - g_p)}{f_p}$$

که در آن :

L_p : کران پایین رده Q_P

g_p : فراوانی تجمعی رده قبل از رده Q_P

f_p : فراوانی رده Q_P

ω : طول رده QP

۱۱- جدول زیر تعداد کتب فروخته شده توسط کتابفروشی را در طول ۳۰ روز نمایش می‌دهد

۱۵	۱۰	۷	۲۰	۱۱	۱۳	۱۸	۶	۵	۴
۱۱	۱۹	۱۲	۱۶	۹	۱۰	۲۱	۱۳	۸	۱۴
۱۷	۲۰	۱۰	۱۲	۱۶	۱۳	۱۱	۱۲	۷	۱۱

برای داده‌های فوق :

الف: یک جدول فراوانی تشکیل دهید و نمودار میله‌ای داده‌ها را رسم کنید.

ب: میانگین داده‌ها $\bar{X}$ ، مد M و میانه m را بدست آورید.

ج: دامنه داده‌ها R ، میانگین انحرافات، واریانس S^2 و ضریب تغییر را بدست آورید.

د: اگر به بزرگترین داده مقدار x $[x \geq 0]$ واحد اضافه کنیم کدامیک از مقادیر مد یا میانه بدون تغییر باقی می‌مانند.

ه: چارک اول و سوم را بدست آورید و دامنه چارکها را محاسبه کنید. (دامنه چارکها : $Q_3 - Q_1$)

و: مقدار دهک چهارم را محاسبه کنید. (دهک چهارم همان $Q_{0.4}$ می‌باشد)

۱۲- در یک شهر میزان درجه حرارت در طول ۳۰ روز به قرار زیر است:

۵	۷	۸/۳	۱۰	۱۱	۱۲/۵	۱۳/۸	۱۳	۱۲	۱۳/۱
۱۴	۱۵/۲	۱۵/۶	۱۶	۱۵/۴	۱۵	۱۶/۵	۱۷	۱۹	۱۹/۷
۲۰/۶	۲۱	۲۱/۳	۲۰/۵	۲۲	۲۲/۸	۲۱/۷	۲۳	۲۴/۱	۲۵

الف: برای داده‌های فوق یک جدول فراوانی تشکیل دهید و هستیوگرام و چند بر فراوانی را رسم کنید

ب: مقادیر میانگین، میانه و مد را محاسبه کنید.

ج: مقادیر واریانس و ضریب تغییر را محاسبه کنید.

د: چند درصد داده‌ها در فاصله $(\bar{X} - S, \bar{X} + S)$ و چند درصد داده‌ها در فاصله $(\bar{X} - 2S, \bar{X} + 2S)$ قرار دارند؟

ه: چند درصد داده‌ها بین چارک اول و سوم قرار دارند؟

فصل دوم

آنالیز ترکیبی و احتمال

در این فصل ابتدا با روشها و قواعد شمارش آشنا می‌شوید. این قوانین به شما کمک می‌کنند که تعداد دفعات قوع یک حالت خاص از بین متممی حالات را به دست بیاورید. در ادامه با استفاده از نتایج بدت آمده و بکارگیری قوانین احتمالات، می‌توانید احتمال وقع یک حالت خاص را بدست آورید.

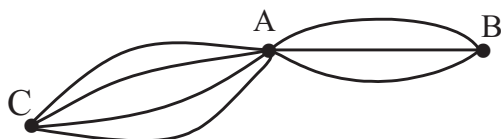
۱-۲ آنالیز ترکیبی

یک سکه را در نظر بگیرید با پرتاب این سکه دو حالت ممکن است رخ بدهد. شیر یا خط حال اگر بخواهیم بدانیم با پرتاب ۳ عدد سکه چند حالت رخ می‌دهد می‌بایستی از قوانین شمارش استفاده کنیم.

به طور کلی تمامی قوانین شمارش مبتنی بر **دو اصل ضرب** و **جمع** می‌باشند. که با استفاده از این دو اصل می‌توان تمامی مسایل شمارش را به سادگی حل نمود.

اصل جمع: اگر کاری را بتوان به n_1 طریق و کار دیگری را به n_2 طریق انجام داد و این دو کار را نتوان همزمان انجام داد، آنگاه کار اول یا کار دوم را می‌توان به $n_1 + n_2$ طریق انجام داد.

مثال ۱: فرض کنیم از شهر A بتوان به سه طریق به شهر B و به چهار طریق به شهر C سفر نمود. به چند طریق می‌توان از شهر A به شهر B یا شهر C سفر نمود؟

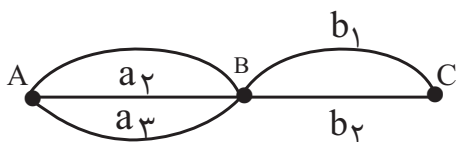


حل: کار اول را سفر از شهر A به B در نظر بگیرید که به $n_1 = 3$ طریق قابل انجام است همینطور کار دوم را سفر از شهر A به C در نظر بگیرید که به $n_2 = 4$ طریق قابل انجام است به وضوح هر دو کار نیز همزمان قابل انجام نیستند پس به $n = n_1 + n_2 = 7$ طریق می‌توان از شهر A به B یا شهر C سفر نمود.

توجه: اصل جمع را می‌توان به n کار نیز تعمیم داد. به شرطی که هیچ دو جفت کاری را نتوان همزمان انجام داد.

اصل ضرب: اگر کاری را بتوان به n_1 طریق و کار دیگری را به n_2 طریق انجام داد و این دو کار را بتوان بصورت همزمان و یکی پس از دیگری انجام داد، آنگاه هر دو کار را می‌توان به $n_1 \times n_2$ طریق انجام داد.

مثال ۱: فرض کنید از شهر A به سه طریق بتوان به شهر B سفر نمود. و از شهر B نیز به دو طریق به شهر C سفر نمود. به چند طریق می‌توان از شهر C سفر نمود؟



با توجه به شکل می‌بایستی ابتدا از شهر A به B و سپس به C سفر نمود که این دو عمل می‌بایستی ابتدا از شهر A به B و سپس به C سفر نمود که این دو عمل می‌بایستی یکی پس از دیگری انجام شوند تا کل کار (سفر از شهر A به C) انجام پذیرد. بنابر این کار اول سفر از شهر A به B که به $n_1 = 3$ طریق قابل انجام است و کار دوم سفر از شهر B به C که به $n_2 = 2$ طریق قابل انجام است و کل کار را می‌توان به $n = n_1 \times n_2 = 3 \times 2 = 6$ طریق می‌توان انجام داد.

اگر راهها از شهر A به B را با a_i و از شهر B به C را با b_i نامگذاری کنیم این ۶ طریق را می‌توان به صورت زیر لیست نمود:

- | | | |
|--------------|--------------|--------------|
| ۱) $a_1 b_1$ | ۳) $a_2 b_1$ | ۵) $a_3 b_1$ |
| ۲) $a_1 b_2$ | ۴) $a_2 b_2$ | ۶) $a_3 b_2$ |

در هر یک از مکانها می‌توانیم هر یک از سه کتاب را قرار دهیم. در مکان اول یکی از سه کتاب را قرار دهیم از آنجا که در هر صورت یک کتاب در مکان اول قرار می‌گیرد در مکان دوم می‌توان یک، از دو کتاب باقیمانده را قرار داد و به همین ترتیب در مکان آخر تنها کتاب باقیمانده قرار می‌گیرد.

از آنجا که این کار را می‌توان در سه مرحله و به صورت پیاپی انجام داد بطوریکه مرحله اول ۳ حالت دوم ۲ حالت و مرحله سوم ۱ حالت دارد. بنابراین با توجه به اصل ضرب داریم:

$$۳ \times ۲ \times ۱ = ۶$$

که فضای نمونه نیز عبارتست از:

$$S = \{(A_1, A_2, A_3), (A_1, A_3, A_2), (A_2, A_1, A_3), (A_2, A_3, A_1), (A_3, A_1, A_2), (A_3, A_2, A_1)\}$$

در حالت کلی تعداد جایگشت‌های n شیئی متمایز در یک ردیف (رعایت ترتیب) عبارتست از:

$$n \times (n-1) \times (n-2) \cdots (2)(1) = n!$$

در نتیجه مثال قبل را می‌توان به صورت خلاصه جایگشت سه شیئی در یک ردیف در نظر گرفت که می‌شود:

$$۳! = ۶$$

توجه: اگر در محاسبه جایگشت‌ها تکرار اشیاء مجاز باشد، می‌توانیم در تمام مکان‌ها از تمام n شیئی استفاده کنیم و در نتیجه تعداد کل جایگشت‌ها عبارتست از: $n \times n \times n \cdots n = n^n$.

مثال ۷: اگر در مثال ۶ از هر سه کتاب به تعداد دلخواه داشته باشیم تعداد کل جایگشت‌ها عبارتست از: $۳ \times ۳ \times ۳ = ۳^۳ = ۲۷$

۲-۱-۲ ترتیب‌های n شیئی به گروه‌های r تایی (P_r^n)

یک حالت کلی‌تر از جایگشت‌ها قرار دادن n شیئی متمایز در r ($r \leq n$) مکان، به ترتیب می‌باشد که با توجه به روش بدست آوردن تعداد کل جایگشت‌ها که گفته شد در این حالت نیز تعداد کل جایگشت‌های n شیئی متمایز در r مکان عبارتست از:

$$n(n-1)(n-2) \cdots (n-r+1) = \frac{n!}{(n-r)!}$$

که آنرا با (P_r^n) می‌دهیم. توجه کنید که $P_n^n = n!$ که همان تعداد کل جایگشت‌هاست. و اگر تکرار اشیاء مجاز باشد در این صورت تعداد حالات ممکن عبارتست از:

$$\underbrace{n \times n \times \cdots \times n}_r = n^r$$

مثال ۸: به چند طریق می‌توان با حروف A, B, C, D, E یک کلمه سه حرفی ساخت؟

$$P_3^5 = \frac{5!}{(5-3)!} = \frac{5!}{2!} = 5 \times 4 \times 3 = ۶۰$$

توجه کنید که حروف متمایز می‌باشند.

مثال ۹: از بین ۱۵ تیم شرکت کننده در مسابقه فوتبال به چند طریق سه تیم رتبه‌های اول و دوم و سوم را بدست می‌آورند؟

$$P_3^{15} = \frac{15!}{(15-3)!} = \frac{15!}{12!} = 15 \times 14 \times 13 = ۲۷۳۰$$

مثال ۱۰: اگر در مثال ۸ تکرار حروف مجاز باشند تعداد حالات ممکن چند تا است؟

$$n^r: ۵ \times ۵ \times ۵ = ۵^۳ = ۱۲۵$$

حال حالت خاصی از جایگشت را در نظر می‌گیریم که در آن تمامی اشیاء متمایز نیستند.

مثال ۱۱: تعداد جایگشت‌های حروف کلمه book را بدست آورید.

در این حالت حرف O دوبار تکرار شده است بنابراین تمامی اشیاء متمایز نیستند اما مجدداً با استفاده از اصل ضرب می‌توانیم تعداد حالات را محاسبه کنیم.

ابتدا تعداد کل حالات را بدون در نظر گرفتن عدم تمایز اشیاء بدست می‌آوریم که برابر است با:

$$n! = 4! = 24$$

دقت کنید که در فضای نمونه مثلاً دو جایگشت obok و obok دوبار تکرار می‌شوند.

زیرا فرض شده است تمام اشیاء (حروف) متمایز باشند اما چون حرف O دوبار تکرار شده است مسلماً با جابجایی دو حرف O در تمام حالات تغییری در کلمه بوجود نمی‌آید بنابراین طبق اصل ضرب $2!$ اضافه در عدد $4!$ ضرب شده است که برای حذف آن می‌بایستی $4!$ را بر $2!$ تقسیم کنیم بنابراین تعداد کل حالات عبارتست از:

$$\frac{4!}{2!} = \frac{24}{2} = 12$$

حالت کلی این مساله بصورت زیر است:

اگر از n شئی، n_1 تای آنها مشابه یکدیگر و n_2 تای آنها مشابه یکدیگر و و n_k تای آنها مشابه یکدیگر باشند تعداد حالاتی که می‌توان این n شی را در یک ردیف مرتب کرد عبارتست از:

$$\frac{n!}{n_1! n_2! \dots n_k!}, \quad \sum_{i=1}^k n_i = n :$$

که آنرا با $(n_1, n_2, \dots, n_k)$ یا $c_n^{n_1, n_2, \dots, n_k}$ نمایش دهیم.

مثال ۱۲: به چند طریق می‌توان دانشجویان یک کلاس ۲۰ نفری را به دسته‌های ۳ و ۴ و ۶ و ۷ نفری تقسیم نمود؟

در این مثال با وجود اینکه تمام اشیاء (دانشجویان) متمایز هستند اما چون هدف تقسیم آنها به چهار دسته ۳ و ۴ و ۶ و ۷ نفری است در نتیجه دانشجویان تخصیص داده شده به هر دسته مثل اشیاء مشابه در نظر گرفته می‌شوند به عبارت بهتر بین تخصیص دانشجوی $\bar{A}$ ام به دسته $\bar{A}$ ام تفاوت قایل هستیم (رعایت ترتیب) اما بین دانشجویان تخصیص داده شده به یک دسته مشخص تفاوتی قایل نیستیم (تشابه) بنابراین تعداد حالات ممکن عبارتست از:

$$\frac{20!}{3! 4! 6! 7!}$$

توجه کنید که می‌بایستی حتماً مجموع تعداد اعضای دسته‌ها برابر تعداد کل دانشجویان باشد»

$$3 + 4 + 6 + 7 = 20$$

۲-۱-۳ ترکیب و مسایل انتخاب

تا بحال با اشیاء متمایز و نظو در ترتیب قرارگیری سروکار داشتیم اما اگر بخواهیم از ترتیب قرار گیری صرف نظر کنیم این حالت به مساله انتخاب r شئی از بین n شئی متمایز تبدیل می‌شود که به آن ترکیب r شئی می‌گوییم.

مثال ۱۳: از بین ۵ مدل اتومبیل سه مدل را انتخاب کنیم به چند طریق این امر امکان‌پذیر است؟

اتومبیل‌ها را از ۱ تا ۵ شماره‌گذاری می‌کنیم و سه جایگاه برای سه مدل اتومبیل انتخابی در نظر می‌گیریم بنابراین طبق جایگشت‌های ۳ شئی از ۵ شئی تعداد کل حالات عبارتست از:

$$P_3^5 = \frac{5!}{(5-3)!} = 60$$

اما چون مساله تنها انتخاب می‌باشد و ترتیب قرار گیری مورد نظر نیست بنابراین به عنوان مثال حالت‌های زیر تکراری محسوب می‌شوند:

$$\begin{array}{cccc} (1, 2, 3) & (3, 2, 1) & (2, 1, 3) & (2, 3, 1) \\ (1, 3, 2) & (3, 1, 2) & & \end{array}$$

به تعداد جایگشت‌های مکان‌های سه اتومبیل، جایگشت تکراری داریم که طبق اصل ضرب یعنی عدد $3!$ اضافه در عدد P_3^5 ضرب شده است و در نتیجه تعداد کل انتخاب‌ها عبارتست از:

$$\frac{P_3^5}{3!} = \frac{5!}{3!2!} = 10$$

در حالت کلی تعداد حالات انتخاب r شیئی از بین n شیئی متمایز عبارتست از:

$$c_n^r = \binom{n}{r} = \frac{P_r^n}{r!} = \frac{n!}{r!(n-r)!} \quad (0 \leq r \leq n)$$

که آنرا با $c(n, r)$ یا $\binom{n}{r}$ نمایش می‌دهند.

در مثال بعد حالتی را بررسی می‌کنیم که بتوان از اشیاء تکراری نیز استفاده نمود:

مثال ۱۴: می‌خواهیم از سه دسته اسکناس ۱۰۰ و ۲۰۰ و ۵۰۰ تومانی ۵ عدد اسکناس انتخاب کنیم، این کار به چند طریق ممکن است؟ همانطور که در این مساله مشاهده می‌کنید اشیاء اسکناس ۱۰۰ و ۲۰۰ و ۵۰۰ تومانی می‌باشند که متمایز می‌باشند. اما می‌توانیم از آنها به تعداد دلخواه انتخاب کنیم. مثلاً حالت‌های زیر ممکن است اتفاق بیفتد:

AAAABB یا ABCBC یا BBBBB

A = ۱۰۰ تومانی

B = ۲۰۰ تومانی

C = ۵۰۰ تومانی

توجه: می‌توانیم از یک اسکناس اصلاً انتخاب نکنیم مثل حالت BBBBB که از اسکناس ۱۰۰ و ۵۰۰ تومانی استفاده نشده است. برای حل مساله حالت زیر را در نظر بگیرید:

ABCBC $\rightarrow$ ABBCC $\rightarrow$ X|XX|XX

زیرا ترتیب مهم نیست

با توجه به حالت بالا A را قبل از پرانتز اول و B را بین دو پرانتز و نهایتاً C را بعد از پرانتز آخر می‌آوریم بنابراین تمام حالت‌ها را می‌توان با قرار دادن دو پرانتز نمایش داد.

پس تعداد کل حالات ممکن از قرار دادن ۵ حرف X و دو پرانتز بدست می‌آید که برابر است با:

$$\binom{3+5-1}{5} = \binom{7}{5} = \frac{7!}{5!2!}$$

در حالت کلی انتخاب r شیئی از بین n شیئی متمایز با تکرار برابر است با:

$$\binom{n+r-1}{r}$$

۲-۲ فضای نمونه و پیشامد

مجموعه تمام حالت‌های ممکن را از انجام یک آزمایش تصادفی فضای نمونه نامیدیم و آنرا با S نمایش داده‌ایم حال به معرفی پیشامد می‌پردازیم. تعریف: هر یک از زیر مجموعه‌های فضای نمونه را یک پیشامد می‌نامیم و در صورتی که نتیجه آزمایش تصادفی عضوی از پیشامد E باشد می‌گوییم پیشامد E به وقوع پیوسته است.

مثال ۱۵: تعداد اعضای فضای نمونه و تعداد پیشامدهای ممکن به ازای پرتاب یک عدد تاس را بدست آورید:

حل: پرتاب یک عدد تاس منجر به شش حالت ۱ تا ۶ می‌شود که در نتیجه فضای نمونه عبارتست از $S = \{1, 2, 3, 4, 5, 6\}$ که ۶ عضو دارد اما چون هر کدام از زیر مجموعه‌های فضای نمونه می‌تواند به عنوان یک پیشامد در نظر گرفته شوند و تعداد کل زیر مجموعه‌های یک مجموعه 2^N ($n =$ تعداد اعضای مجموعه) می‌باشد بنابراین تعداد پیشامدهای ممکن $2^6 = 64$ می‌باشد.

به عنوان مثال هر کدام از مجموعه‌های A, B, C, D پیشامدهایی برای فضای نمونه S می‌باشند:

$$A = \{\phi\} \quad D = \{1, 2, 3\} \quad C = \{2\} \quad D = \{1, 2, 3, 4, 5, 6\}$$

۲-۱ انواع پیشامدها و اعمال روی پیشامدها

با توجه به تعریف پیشامد حالت‌های مختلفی برای پیشامدها خواهیم داشت که عبارتند از:

پیشامد ساده: پیشامدی که تنها دارای یک عضو باشد را پیشامد ساده می‌نامیم مثل

$$E = \{\text{خط و شیر}\}$$

که در نتیجه حاصل از پرتاب دو سکه با هم می‌باشد.

$$E = \{1, 2\}$$

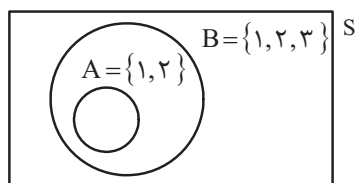
پیشامد مرکب: پیشامدی دارای دو عضو یا بیشتر را مرکب می‌نامیم مثل:

پیشامد تهی یا محال: اگر پیشامدی دارای هیچ عضوی نباشد به آن پیشامد محال گویند.

پیشامد حتمی: پیشامدی که برابر با فضای نمونه S باشد پیشامد حتمی نامیده می‌شود مثلاً در پرتاب یک عدد تاس $E = \{1, 2, 3, 4, 5, 6\}$ پیشامد حتمی است.

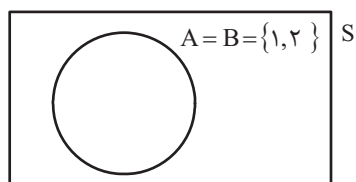
از آنجا که پیشامدها خود یک مجموعه می‌باشند می‌توان اعمالی که روی مجموعه‌ها تعریف می‌شود را روی پیشامدها نیز اعمال نمود. دو پیشامد A و B فضای نمونه را S در نظر می‌گیریم.

۱- هر گاه وقوع پیشامد A وقوع پیشامد B را نتیجه دهد گوئیم پیشامد A زیر پیشامد B می‌باشد و آنرا بصورت $A \subset B$ نمایش می‌دهیم. این مطلب را در شکل نیز مشاهده می‌کنید.



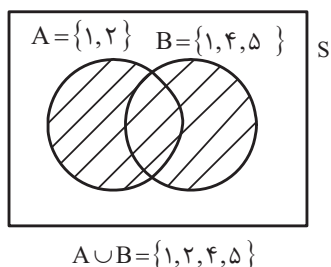
۲- هر گاه وقوع پیشامد A وقوع پیشامد B را نتیجه دهد و بالعکس دو پیشامد را مساوی گویند.

$$A = B \Leftrightarrow (A \subset B, B \subset A)$$



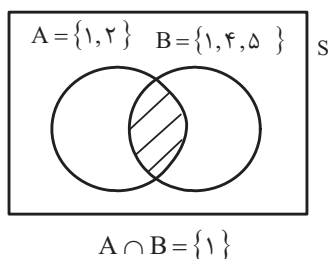
۳- اجتماع دو پیشامد: در صورتی که حداقل یکی از دو پیشامد A و B رخ دهند گویند اجتماع دو پیشامد یا $A \cup B$ رخ داده است و آنرا بصورت زیر می‌نویسیم:

$$A \cup B = \{x \mid x \in A \text{ یا } x \in B\}$$



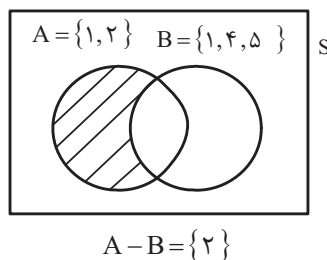
۴- اشتراک دو پیشامد: در صورتی که هر دو پیشامد A و B بصورت همزمان رخ دهند گویند اشتراک آندو یا $A \cap B$ رخ داده است.

$$A \cap B = \{x | x \in A, x \in B\}$$



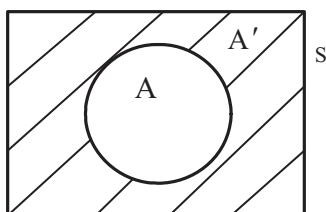
۵- تفاضل دو پیشامد: اگر فقط A رخ دهد در حالی که B رخ نداده باشد گویند تفاضل دو پیشامد یا $A - B$ رخ داده است.

$$A - B = \{x | x \in A, x \notin B\}$$

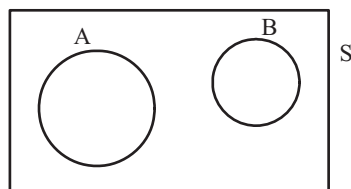


۶- متمم یک پیشامد: A' را متمم پیشامد A در نظر می‌گیریم وقتی که A' به معنی عدم وقوع A می‌باشد.

$$A' = \{x | x \in S, x \notin A\}$$



۷- در صورتی که دو پیشامد همزمان نتواند رخ دهند یا به عبارتی اشتراکی نداشته باشند آندو را دو پیشامد ناسازگار گوئیم در این حالت $A \cap B = \emptyset$



۳-۲ احتمال

به هر یک از پیشامدها می‌توان عددی نسبت داد که معرف میزان احتمال وقوع آن پیشامد باشد. برای این منظور می‌بایستی از قواعدی پیروی کنیم که برای تمام پیشامدها یکی باشد با استفاده از تعریف فراوانی نسبی می‌توان به محاسبه احتمال پرداخت به این ترتیب که یک آزمایش را N مرتبه انجام می‌دهیم و تعداد دفعاتی که پیشامد A مشاهده شده را $n(A)$ در نظر می‌گیریم در این صورت فراوانی نسبی وقوع پیشامد A عبارتست از $\frac{n(A)}{N}$ که آنرا احتمال وقوع پیشامد A نامیده و با $P(A)$ نمایش می‌دهیم.

فراوانی نسبی همواره عددی بین صفر و یک می‌باشد بنابراین مقدار احتمال نیز در این بازه می‌باشد از طرفی احتمال وقوع تمام اعضای فضای نمونه برابر ۱ می‌باشد یعنی $p(S) = 1$ از طرفی تعداد دفعات مشاهده یکی از دو پیشامد A یا B در صورتی که آندو با یکدیگر اشتراکی نداشته باشند برابر است با تعداد دفعات مشاهده $A \cap B$ بنابراین $p(A \cup B) = p(A) + p(B)$ به شرطی که $A \cap B = \emptyset$ سه قلعه بدست آمده در بالا، اصول موضوعه احتمال می‌گوئیم و بطور کلی احتمالات باید تابع قوانین زیر باشند:

$$1 - p(S) = 1$$

$$0 \leq p(A) \leq 1 \quad \forall A \subset S$$

$$p(A_1 \cup A_2 \cup \dots) = p(A_1) + p(A_2) + \dots$$

$$A_i \cap A_j = \emptyset, \quad \forall i \neq j$$

۱.۳.۲ قوانین احتمال

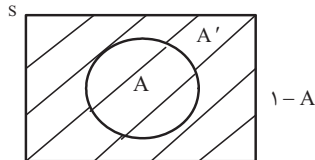
قبلاً با استفاده از قوانین مجموعه‌ها اعمالی را روی پیشامدها تعریف کردیم حالا با استفاده از آنها و اصول موضوعه احتمال، چندین قانون برای احتمالات بدست می‌آوریم.

قضیه: $p(\phi) = 0$

اثبات: $S \cup \phi = S$ پس بنابر اصل ۳

$$p(S \cup \phi) = p(S) + p(\phi) = 1 + p(\phi) \Rightarrow 1 + p(\phi) = 1 \Rightarrow p(\phi) = 0$$

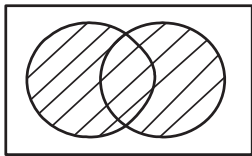
با استدلالی مشابه می‌توان قضایای زیر را اثبات نمود:



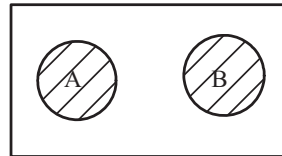
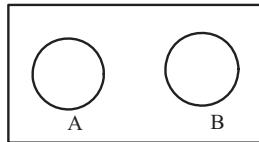
$$p(A') = 1 - p(A) \quad -1$$

$$p(A \cup B) = p(A) + p(B) - p(A \cap B) \quad -2$$

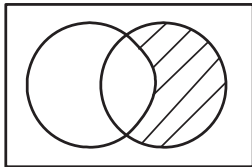
۳- اگر A و B ناسازگار باشند داریم $A \cap B = \phi$ و در نتیجه



$$p(A \cup B) = 0$$



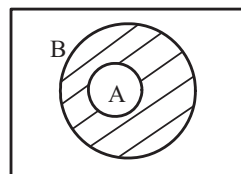
$$p(A \cup B) = p(A) + p(B)$$



$$p(A \cup B) = p(B \cap A') = p(B) - p(A \cap B) \quad -4$$

۵- اگر $A \subset B$ انگاه:

$$\begin{cases} p(B - A) = p(B) - p(A) \\ p(A) \leq p(B) \end{cases}$$



مثال ۱۶: سه سکه را همزمان پرتاب می‌کنیم احتمال بدست آوردن سه شیر یا سه خط چقدر است؟

حل: پیشامد A را بدست آوردن سه شیر در نظر می‌گیریم و به همین ترتیب پیشامد B را بدست آمدن سه خط

$$A = \{\text{ش و ش و ش}\} \quad B = \{\text{خ و خ و خ}\}$$

از آنجا که $A \cap B = \phi$ بنابراین A و B ناسازگار هستند و در نتیجه پیشامد بدست آمدن سه شیر یا سه خط برابر است با:

$$p(A \cup B) = p(A) + p(B)$$

با توجه به قواعد شمارش تعداد اعضای فضای نمونه برابر است با: $n(S) = 2^3 = 8$, $n(A) = 1$

$$p(A) = \frac{n(A)}{n(S)} = \frac{1}{8}$$

$$p(B) = \frac{n(B)}{n(S)} = \frac{1}{8}$$

$$p(A \cup B) = p(A) + p(B) = \frac{1}{8} + \frac{1}{8} = \frac{2}{8}$$

بنابراین:

پس:

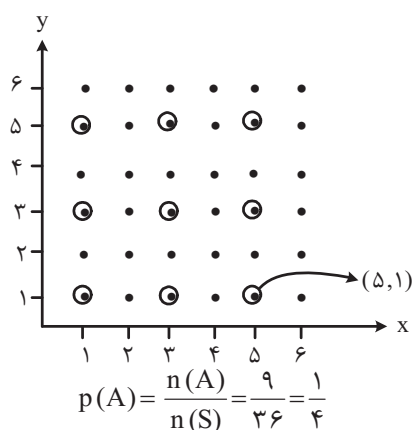
مثال ۱۷: یک جفت تاس را می‌ریزیم احتمال بدست آمدن اعدادی فرد روی هر دو تاس چقدر است؟

$$A = \{\text{آمدن اعداد فرد بر روی دو تاس}\}$$

$$n(A) = 3 \times 3 = 9$$

سه عدد فرد بر روی تاس اول $\{1, 3, 5\}$

در شکل روبرو اعضای مجموعه A در دوایری نشان داده شده‌اند:



$$n(S) = 6 \times 6 = 36 \text{ از تعداد کل حالات}$$

مثال ۱۸: یک سکه را ۵ مرتبه پرتاب می‌کنیم احتمال اینکه حداقل دو بار شیر بیاید چقدر است؟

A = پیشامد اینکه حداقل ۲ بار شیر بیاید

$$A' = \text{پیشامد یا اصلاً شیر نیاید یا دقیقاً یکبار شیر بیاید} = \left\{ \begin{array}{l} (\text{خ و خ و خ و خ و خ و خ}) \text{ و } (\text{خ و خ و خ و خ و ش}) \text{ و } (\text{خ و خ و خ و ش و خ}) \text{ و } (\text{خ و خ و ش و خ و خ}) \\ (\text{ش و خ و خ و خ و خ}) \text{ و } (\text{خ و ش و خ و خ و خ}) \text{ و } (\text{خ و خ و ش و خ و خ}) \end{array} \right\}$$

$$n(S) = 2^5 = 32$$

$$p(A') = \frac{6}{32} = \frac{3}{16}$$

$$p(A) = 1 - p(A') = 1 - \frac{3}{16} = \frac{13}{16}$$

مثال ۱۹: الف) تعداد جوابهای صحیح و غیر منفی معادله $x_1 + x_2 + \dots + x_n = r$ را در صورتیکه $x_i \geq 0, i=1, 2, \dots, n$ باشد را بدست آورید؟

ب) می‌خواهیم ۷ عدد بلیط را بین سه نفر تقسیم کنیم بطوریکه به هر نفر حداقل یک بلیط برسد مطلوبست احتمال اینکه به نفر دوم حداقل ۲ بلیط برسد؟

حل: الف) این مساله مشابه حالتیست که بخواهیم r شئی را بین n نفر تقسیم کنیم بطوریکه در این حالت به هر شخص می‌توان بیشتر از یک شئی داد و ترتیب اهمیتی ندارد در قسمت قوانین شمارش نشان دادیم که جواب چنین مساله‌ای با استفاده از رابطه $\binom{n+r-1}{r}$ بدست می‌آید.

ب) تعداد کل حالات فضای نمونه در این حالت از حل معادله $x_1 + x_2 + x_3 = 7, x_i \geq 1$ در نظر گرفته شود اما برای استفاده از رابطه $\binom{n+r-1}{r}$ باید $x_i \geq 0$ باشد بنابراین از تغییر متغیر زیر استفاده می‌کنیم.

$$x_i \geq 1 \rightarrow x_i - 1 \geq 0 \quad y_i = x_i - 1 \Rightarrow y_i \geq 0$$

$$\Rightarrow x_i = y_i + 1 \quad y_1 + 1 + y_2 + 1 + y_3 + 1 = 7$$

$$\Rightarrow y_1 + y_2 + y_3 = 7 - 3 = 4$$

بنابراین مساله به صورت زیر خلاصه می شود:

$$\begin{cases} y_1 + y_2 + y_3 = 4 \\ y_i \geq 0 \end{cases}$$

$$n(s) = \binom{3+4-1}{4} = 15$$

که جواب آن عبارتست از ۱۵

برای بدست آوردن تعداد حالات مطلوب باید مساله زیر را حل کنیم:

$$\begin{cases} x_1 + x_2 + x_3 = 7 \\ x_1, x_3 \geq 1 \\ x_2 \geq 2 \end{cases}$$

دوباره با استفاده از تغییر متغیر داریم:

$$y_1 = x_1 - 1, \quad y_3 = x_3 - 1, \quad y_2 = x_2 - 2$$

$$\Rightarrow y_i \geq 0$$

که در نتیجه مساله به فرم زیر تبدیل می شود:

$$A: \begin{cases} y_1 + y_2 + y_3 = 3 \\ y_i \geq 0 \end{cases}$$

$$n(A) = \binom{3+3-1}{3} = 10$$

که جواب آن عبارتست از:

بنابراین احتمال فوق عبارتست از:

$$p(A) = \frac{n(A)}{n(S)} = \frac{10}{15} = \frac{2}{3}$$

۴.۲ احتمال روی فضای نمونه نامتناهی

با توجه به نوع مسایل فضای نمونه می تواند نامتناهی باشد در این حالت مجموعه پیشامدها نیز نامتناهی خواهند بود البته احتمال پیشامدها می بایستی

$$\sum_{i=1}^{+\infty} p(A_i) = 1, \quad 0 \leq p(A_i) \leq 1$$

در اصول موضوعه صدق کند بطوریکه

مثال ۲۰: A و B به ترتیب به سوی هدفی شلیک می کنند A با احتمال $\frac{1}{4}$ و B با احتمال $\frac{3}{4}$ هدف را مورد اصابت قرار می دهند مطلوبست احتمال اینکه B زودتر از A هدف را بزند؟ (در صورتی که A اول شروع کند).

$$p(A') = 1 - \frac{1}{4} = \frac{3}{4}$$

پیشامد اینکه A هدف را بزند: A'

احتمال اینکه B به هدف نزند: B'

برای اینکه B زودتر هدف را مورد اصابت قرار دهد حالات زیر را خواهیم داشت:

پیشامد	$A'B$	$A'B'A'B$	$A'B'A'B'A'B'$	...
احتمال	$\frac{1}{4} \times \frac{3}{4}$	$(\frac{1}{4})^2 \times \frac{3}{4} \times \frac{1}{4}$	$(\frac{1}{4})^3 (\frac{1}{4})^2 \times \frac{3}{4}$	...

$$p(B') = 1 - \frac{3}{4} = \frac{1}{4}$$

از آنجا که پیشامدها ناسازگار هستند باید احتمال اجتماع تمام پیشامدها محاسبه شود که برابرست با مجموع احتمالات هر یک از پیشامدها یعنی:

$$p(A'B \cup A'B'A'B \cup \dots) = p(A'B) + p(A'B'A'B) + \dots = \frac{1}{2} \times \frac{3}{4} + \left(\frac{1}{2}\right)^2 \times \frac{3}{4} + \dots$$

$$= \frac{3}{4} \times \frac{1}{2} \left(1 + \frac{1}{2} + \left(\frac{1}{2}\right)^2 + \dots\right) = \frac{3}{4} \times \frac{1}{1 - \frac{1}{2}} = \frac{3}{2}$$

مثال ۲۱: سکه‌ای را آنقدر پرتاب می‌کنیم تا برای اولین بار شیر مشاهده شود مطلوبست:

الف) احتمال آنکه تعداد زوجی پرتاب لازم باشد.

ب) احتمال آنکه حداکثر ۱۰ پرتاب لازم باشد.

حل: برای اینکه برای بار اول شیر مشاهده کنیم می‌بایستی در تمام پرتاب‌های قبلی خط مشاهده شده باشد در این صورت احتمال‌ها را می‌توان در جدول زیر خلاصه کنیم:

S	e_1 ش	e_2 ش خ	e_3 ش خ خ	e_4 ش خ خ خ	...
P	$\frac{1}{2}$	$\frac{1}{4}$	$\frac{1}{8}$	$\frac{1}{16}$	...

پیشامد اینکه تعداد زوجی پرتاب لازم باشد: A

الف) برای اینکه تعداد زوجی پرتاب لازم باشد باید پیشامدهای $e_2, e_4, e_6, \dots$ رخ دهند بنابراین:

$$p(A) = \frac{1}{4} + \frac{1}{16} + \frac{1}{64} + \dots = \frac{\frac{1}{4}}{1 - \frac{1}{4}} = \frac{1}{3}$$

همچنین می‌توان احتمال اینکه تعداد فردی پرتاب لازم باشد را از روی $p(A)$ محاسبه نمود. چون پیشامد تعداد فردی پرتاب متمم پیشامد A می‌باشد بنابراین:

$$p(A') = 1 - p(A) = 1 - \frac{1}{3} = \frac{2}{3}$$

ب) پیشامد آنکه حداکثر ۱۰ پرتاب لازم باشد $B =$

برای محاسبه احتمال پیشامد B می‌بایستی مجموع ۱۰ جمله اول از جدول را بدست بیاوریم اما با استفاده از متمم این پیشامد می‌توان تعداد آنرا به راحتی محاسبه نمود:

پیشامد اینکه حداقل ۱۱ پرتاب لازم باشد $B' =$

$$p(B') = \left(\frac{1}{2}\right)^{11} + \left(\frac{1}{2}\right)^{12} + \dots = \frac{\left(\frac{1}{2}\right)^{11}}{1 - \frac{1}{2}} = \left(\frac{1}{2}\right)^{10}$$

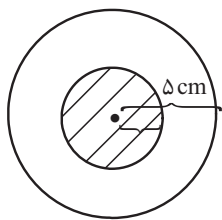
$$p(B) = 1 - p(B') = 1 - \left(\frac{1}{2}\right)^{10} = \frac{1023}{1024} \approx 0.999$$

۵.۲ فضای نمونه پیوسته

در شرایطی می‌توان فضای نمونه را طوری تعریف نمود که بصورت یک بازه یا یک سطح در نظر گرفته شود. مثلاً $S = [1, 2]$ در این حالت هر زیر بازه مثل $E = [1/5, 1/75]$ می‌تواند به عنوان یک پیشامد در نظر گرفته شود. برای محاسبه احتمال پیشامد باز هم حالت مطلوب را به کل فضای نمونه در نظر می‌گیریم یعنی:

$$p(A) = \frac{\text{مساحت ناحیه A}}{\text{مساحت فضای نمونه S}} \quad \text{یا} \quad \frac{\text{طول بازه A}}{\text{طول فاصله S}}$$

مثال ۲۲: تیراندازی به هدفی شلیک می‌کند که قطر آن ۱۰ سانتیمتر و قطر دایره مرکزی هدف ۲ سانتیمتر است احتمال اینکه تیر به مرکز هدف اصابت کند چقدر است؟



با توجه به شکل مساحت فضای نمونه عبارتست از: 25π

و مساحت ناحیه مطلوب عبارتست از: π

بنابراین احتمال مورد نظر عبارتست از:

$$p(A) = \frac{\pi}{25\pi} = \frac{1}{25}$$

مثال ۲۳: در مثال قبل احتمال اینکه تیرانداز دقیقاً مرکز هدف را بزند چقدر است؟

از آنجا که مرکز هدف یک نقطه محسوب می‌شود بنابراین مساحت آن صفر واحد می‌باشد و در نتیجه احتمال مورد نظر برابر صفر خواهد بود.

۶.۲ احتمال شرطی

در بسیاری از مواقع می‌دانیم که پیشامد A رخ داده و می‌خواهیم احتمال رخ دادن پیشامد B را مشروط بر اینکه A رخ داده است بدست بیاوریم در این حالت می‌بایستی از تعاریف احتمال شرطی استفاده کنیم. به مثال زیر توجه کنید:

مثال ۲۴: یک تاس طوری طراحی شده است که احتمال آمدن هر عدد متناسب با آن عدد می‌باشد در این صورت:

الف) احتمال رخ دادن یک عدد زوج را بیابید.

ب) اگر بدانیم عدد حاصل از پرتاب کوچکتر یا مساوی با عدد ۴ می‌باشد احتمال آمدن یک عدد فرد را بیابید.

ابتدا احتمال هر عدد را محاسبه می‌کنیم با توجه به اینکه احتمال آمدن هر عدد باید متناسب با آن عدد باشد معادله زیر را داریم:

$$x + 2x + 3x + 4x + 5x + 6x = 1$$

که در آن x احتمال آمدن عدد ۱ می‌باشد.

$$21x = 1 \rightarrow x = \frac{1}{21}$$

بنابراین:

S	۱	۲	۳	۴	۵	۶
احتمال	$\frac{1}{21}$	$\frac{2}{21}$	$\frac{3}{21}$	$\frac{4}{21}$	$\frac{5}{21}$	$\frac{6}{21}$

توجه کنید که مجموع احتمالات برابر ۱ می‌باشد.

$$A = \{2, 4, 6\}$$

الف) پیشامد رخ دادن عدد زوج $A =$

$$p(A) = \frac{2}{21} + \frac{4}{21} + \frac{6}{21} = \frac{12}{21}$$

ب) در این حالت مجموعه فضای نمونه محدود می‌باشد به $B = \{1, 2, 3, 4\}$ زیرا می‌دانیم عدد بدست آمده کوچکتر یا مساوی با ۴ می‌باشد. از طرفی چون احتمال آمدن هر عدد متناسب با همان عدد است پس مدل احتمال روی این فضای نمونه جدید عبارتست از:

$$x + 2x + 3x + 4x = 1 \rightarrow 10x = 1 \rightarrow x = \frac{1}{10}$$

	۱	۲	۳	۴
احتمال	$\frac{1}{10}$	$\frac{2}{10}$	$\frac{3}{10}$	$\frac{4}{10}$

(پیشامد رخ دادن عدد فرد به شرط وقوع B) $C =$

بنابراین احتمال وقوع یک عدد فرد به شرطی که پیشامد B رخ داده باشد عبارتست از:

$$p(C) = \frac{1}{10} + \frac{3}{10} = \frac{4}{10} = \frac{2}{5}$$

در حالت کلی تعریف زیر را خواهیم داشت:

تعریف: احتمال وقوع پیشامد B به شرط وقوع پیشامد A که آنرا با نماد $P(B|A)$ نمایش می‌دهیم به صورت زیر تعریف می‌شود:

$$P(B|A) = \frac{P(A \cap B)}{P(A)}, \quad P(A) \neq 0$$

مثال ۲۵: با استفاده از تعریف احتمال مثال ۲۴ (ب) را محاسبه می‌کنیم.

$A = \{1, 2, 3, 4\}$ پیشامد وقوع عددی کوچکتر از ۴

$B = \{1, 3, 5\}$ پیشامد رخ دادن عددی فرد

$$P(A) = \frac{1}{21} + \frac{2}{21} + \frac{3}{21} + \frac{4}{21} = \frac{10}{21}$$

$$A \cap B = \{1, 3\} \rightarrow P(A \cap B) = \frac{1}{21} + \frac{3}{21} = \frac{4}{21}$$

$$\Rightarrow P(B|A) = \frac{P(A \cap B)}{P(A)} = \frac{\frac{4}{21}}{\frac{10}{21}} = \frac{4}{10} = \frac{2}{5}$$

ملاحظه می‌کنید که در این حالت دیگر نیازی به محاسبه جدول احتمالات میانی نمی‌باشد.

مثال ۲۶: در جعبه‌ای ۵ مهره به رنگ آبی و ۴ مهره به رنگ قرمز موجود می‌باشند از این جعبه ۳ مهره به تصادف خارج می‌کنیم اگر بدانیم دو مهره

از سه مهره به رنگ آبی می‌باشند احتمال اینکه مهره سوم به رنگ قرمز باشد چقدر است؟

$A =$ پیشامد اینکه دو مهره از سه مهره آبی باشند

$B =$ پیشامد اینکه یک مهره از سه مهره قرمز باشد

$A \cap B =$ پیشامد اینکه دو مهره آبی و یک مهره قرمز باشد

$$P(A \cap B) = \frac{\binom{4}{1} \binom{5}{2}}{\binom{9}{3}} \quad P(A) = \frac{\binom{5}{2} \binom{4}{1} + \binom{5}{3}}{\binom{9}{3}}$$

$$P(B|A) = \frac{P(A \cap B)}{P(A)} = \frac{\frac{\binom{5}{2} \binom{4}{1}}{\binom{9}{3}}}{\frac{\binom{5}{2} \binom{4}{1} + \binom{5}{3}}{\binom{9}{3}}}$$

۲. ۷ قانون ضرب احتمال

در صورتی که فرمول $P(A|B) = \frac{P(A \cap B)}{P(B)}$ را ساده کنیم خواهیم داشت:

$$P(A|B) = \frac{P(A \cap B)}{P(B)} \Rightarrow P(A \cap B) = P(A|B) P(B)$$

$$P(B|A) = \frac{P(A \cap B)}{P(A)} \Rightarrow P(A \cap B) = P(B|A) P(A)$$

در نتیجه فرمول کلی زیر را بدست می آوریم:

$$P(A \cap B) = P(A|B)P(B) = P(B|A)P(A)$$

این رابطه در حالتی که دو پیشامد A و B بتوانند بصورت همزمان رخ دهند به کار می رود و به آن قانون ضرب احتمال می گوئیم.

مثال ۲۸: جدول زیر احتمال شاغل بودن مردان و زنان را در یک جامعه آماری نشان می دهد:

مثلاً در جدول احتمال اینکه شخص مرد باشد ۰/۵۷ و احتمال اینکه مرد باشد و شاغل هم باشد ۰/۵۲ می باشد.

		M مرد	F زن	جمع
E	شاغل	۰/۵۲	۰/۴۱	۰/۹۳
U	بیکار	۰/۰۵	۰/۰۲	۰/۰۷
	جمع	۰/۵۷	۰/۴۳	۱/۰۰

با توجه به جدول به سوالات زیر پاسخ دهید:

الف) اگر از جامعه فوق یک نفر انتخاب کنیم و بدانیم شاغل است احتمال اینکه مرد باشد چقدر است؟

ب) احتمال اینکه نفر انتخابی از جامعه فوق شاغل باشد؟

ج) اگر نفر انتخابی مرد باشد احتمال اینکه شاغل باشد چقدر است؟

حل:

الف) تعریف می کنیم:

U = پیشامد اینکه نفر انتخابی بیکار باشد

E = پیشامد اینکه نفر انتخابی شاغل باشد

M = پیشامد اینکه نفر انتخابی مرد باشد

F = پیشامد اینکه نفر انتخابی زن باشد

بنابراین مقدار $P(E|M)$ را می خواهیم:

با توجه به جدول داریم: $P(E \cap M) = ۰/۵۲$

$$P(M) = P(M \cap E) + P(M \cap U) = ۰/۵۲ + ۰/۰۵ = ۰/۵۷$$

$$P(E|M) = \frac{P(E \cap M)}{P(M)} = \frac{۰/۵۲}{۰/۵۷} = ۰/۹۱$$

ب) پس:

$$P(E) = P(E \cap M) + P(E \cap F) = ۰/۵۲ + ۰/۴۱ = ۰/۹۳$$

$$P(M|E) = \frac{P(M \cap E)}{P(E)} = \frac{0.52}{0.93} = 0.56$$

(ج)

مثال ۲۹: فرض کنید در مثال قبل تنها اطلاعات زیر را در اختیار داشته باشیم:

$$P(M) = 0.57$$

$$P(E|M) = 0.91$$

$$P(U|M) = 0.09$$

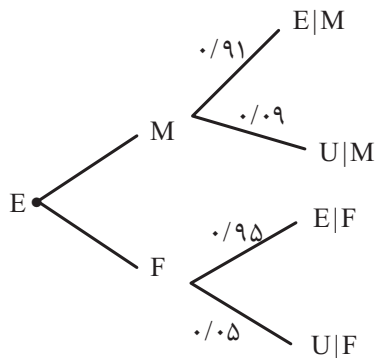
$$P(F) = 0.43$$

$$P(E|F) = 0.95$$

$$P(U|F) = 0.05$$

در اینصورت مقادیر $P(M|E)$ و $P(E)$ را محاسبه کنید.

حل: احتمالات داده شده را می‌توان به صورت درختی در نظر گرفت در این صورت داریم:



که طبق قانون ضرب احتمال:

$$P(E) = P(E \cap M) + P(E \cap F) = P(E|M) P(M) + P(E|F) P(F)$$

$$= 0.91 \times 0.57 + 0.95 \times 0.43 = 0.93$$

$$P(M|E) = \frac{P(M \cap E)}{P(E)} = \frac{P(E|M) \times P(M)}{P(E \cap M) + P(E \cap F)}$$

$$= \frac{P(E|M) P(M)}{P(E|M) P(M) + P(E|F) P(F)} = \frac{0.91 \times 0.57}{0.91 \times 0.57 + 0.95 \times 0.43} = \frac{0.52}{0.93} = 0.56$$

۲. ۸. پیشامدهای مستقل

رابطه $P(A|B) = \frac{P(A \cap B)}{P(B)}$ را در نظر بگیرید، اگر تحت شرایطی مقدار $P(A \cap B)$ برابر $P(A) \times P(B)$ شود خواهیم داشت:

$$P(A|B) = \frac{P(A) \times P(B)}{P(B)} = P(A)$$

به عبارتی احتمال وقوع A به شرط B برابر با احتمال وقوع A می‌باشد و این یعنی دانستن اینکه پیشامد B رخ داده است هیچ اطلاعاتی در مورد رخ دادن پیشامد A بدست نمی‌دهد بنابراین می‌توان گفت دو پیشامد از یکدیگر مستقل هستند و تعریف زیر را خواهیم داشت»
تعریف: دو پیشامد A و B را از یکدیگر مستقل می‌نامیم اگر داشته باشیم:

$$P(A \cap B) = P(A) \times P(B)$$

مثال ۳۰: در یک ایستگاه مترو احتمال اینکه قطار به موقع در ایستگاه قرار داشته باشد 0.9 می‌باشد مطلوبست:

الف) احتمال اینکه قطار سه روز متوالی به موقع در ایستگاه باشد.

ب) احتمال اینکه قطار در روز سوم برای بار دوم دیر کند.

حل:

الف) ابتدا پیشامد زیر را تعریف می‌کنیم:

احتمال اینکه قطار در روز A به موقع در ایستگاه قرار داشته باشد $A =$ در این صورت می‌توان پیشامدهای A_1, A_2, A_3 را از یکدیگر مستقل در نظر گرفت زیرا پیشامد اینکه قطار امروز دیر کند به پیشامد اینکه قطار فردا هم دیر کند ارتباطی ندارد بنابراین با توجه به فرمول استقلال داریم:

$$P(A_1 \cap A_2 \cap A_3) = P(A_1) \times P(A_2) \times P(A_3) = 0.9 \times 0.9 \times 0.9 = 0.729$$

ب) پیشامد زیر را تعریف می‌کنیم:

پیشامد اینکه قطار در روز A_i دیر کند $A'_i =$

$$P(A'_i) = 1 - P(A_i) = 1 - 0.9 = 0.1$$

برای اینکه قطار در روز سوم برای بار دوم دیر کند می‌بایستی در دو روز قبل حداقل یکبار دیر کرده باشد. احتمال زیر را محاسبه می‌کنیم:

$$\begin{aligned} P[(A'_1 \cap A'_2 \cap A'_3) \cup (A_1 \cap A'_2 \cap A'_3)] &= \\ &= P(A'_1 \cap A'_2 \cap A'_3) + P(A_1 \cap A'_2 \cap A'_3) \\ &= P(A'_1) P(A'_2) P(A'_3) + P(A_1) P(A'_2) P(A'_3) \\ &= 0.1 \times 0.1 \times 0.1 + 0.9 \times 0.1 \times 0.1 = 0.118 \end{aligned}$$

استقلال سه پیشامد: سه پیشامد A, B, C را مستقل از یکدیگر می‌نامیم اگر داشته باشیم:

$$P(A \cap B) = P(A) P(B) \quad -1$$

$$P(A \cap C) = P(A) P(C) \quad -2$$

$$P(B \cap C) = P(B) P(C) \quad -3$$

$$P(A \cap B \cap C) = P(A) P(B) P(C) \quad -4$$

پیشامدهای ناسازگار و مستقل: به تفاوت‌های پیشامدهای ناسازگار و مستقل توجه کنید:

- اگر دو پیشامد ناسازگار باشند: نمی‌توانند همزمان رخ دهند، احتمال اشتراک آنها صفر می‌باشد، احتمال اجتماع آنها برابر مجموع احتمالات هر یک می‌باشد.

- اگر دو پیشامد مستقل از یکدیگر باشند: هر دو می‌توانند همزمان رخ دهند، احتمال اشتراک آنها برابر با حاصل ضرب احتمالات آنهاست، احتمال اجتماع آنها کوچکتر یا مساوی با مجموع احتمالات هر یک می‌باشد.

* توجه کنید که اگر دو پیشامد بخواهند بصورت همزمان هم ناسازگار باشند و هم مستقل در نتیجه باید داشته باشیم:

$$\begin{cases} P(A \cap B) = 0 \\ P(A \cap B) = P(A) P(B) \end{cases} \Rightarrow P(A) P(B) = 0 \Rightarrow P(A) = 0 \quad \text{یا} \quad P(B) = 0$$

۹.۲ فرمول احتمال بینر و فرمول تفکیک احتمال

فرمول احتمال بینر روش ساده‌تری برای محاسبه احتمالات شرطی در حالتی که اطلاعات کمی در مورد مساله داریم ارایه می‌کند که در ذیل نحوه بدست آوردن آنرا شرح می‌دهیم:

فرض کنید فضای نمونه S را به دو پیشامد A و A' که متمم آن می‌باشد تقسیم کنیم به صورت شکل زیر:حال برای حل مسئله در حالت کلی پیشامد B را در این فضا طوری در نظر می‌گیریم که با A و A' اشتراک داشته باشد:

می‌خواهیم احتمال A به شرط وقوع B را محاسبه کنیم:

$$P(A|B) = \frac{P(A \cap B)}{P(B)} \quad 1-1$$

از روی شکل می‌توان به راحتی احتمال وقوع B را محاسبه نمود.

$$B = (A \cap B) \cup (A' \cap B) \\ P(B) = P(A \cap B) + P(A' \cap B) = P(A) P(B|A) + P(A') P(B|A') \quad 1-2$$

به فرمول فوق فرمول تفکیک احتمال گویند که حالت کلی‌تر آنرا در ادامه بدست می‌آوریم.

$$P(A \cap B) = P(A) P(B|A) \quad 1-3$$

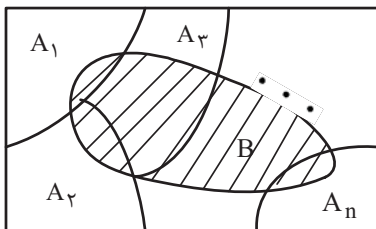
$$1-1, 1-2, 1-3 \Rightarrow P(A|B) = \frac{P(A) P(B|A)}{P(A) P(B|A) + P(A') P(B|A')} \quad 1-4$$

رابطه 1-4 به فرمول بنیر معروف است.

برای بدست آوردن حالت کلی‌تر روابط 1-2 و 1-4 فرض کنید فضای نمونه به پیشامدهای $A_1, A_2, \dots, A_n$ طوری تقسیم شده باشد که به ازای

$$S = \bigcup_{i=1}^n A_i \text{ و } A_i \cap A_j = \emptyset \text{ برای } i \neq j \text{ داشته باشیم}$$

مطابق شکل زیر:



در این حالت اصطلاحاً می‌گوییم فضای نمونه S به n پیشامد A_1 تا A_n افراز شده است. و پیشامد B نیز پیشامد دلخواه در فضای نمونه S باشد در این صورت داریم:

$$B = B \cap S = B \cap \left(\bigcup_{i=1}^n A_i \right) = (B \cap A_1) \cup (B \cap A_2) \cup \dots \cup (B \cap A_n)$$

که در این رابطه پیشامدهای $(B \cap A_1), (B \cap A_2), \dots, (B \cap A_n)$ دو به دو ناسازگارند. بنابراین داریم:

$$P(B) = P(B \cap A_1) + P(B \cap A_2) + \dots + P(B \cap A_n) = P(A_1) P(B|A_1) + P(A_2) P(B|A_2) \\ + \dots + P(A_n) P(B|A_n) = \sum_{i=1}^n P(A_i) P(B|A_i)$$

این رابطه فرمول تفکیک احتمال یا فرمول احتمال کل نامیده می‌شود. به همین ترتیب برای $P(A_i|B)$ داریم:

$$P(A_i|B) = \frac{P(A_i \cap B)}{P(B)} = \frac{P(A_i) P(B \cap A_i)}{\sum_{i=1}^n P(A_i) P(B \cap A_i)}$$

به این رابطه فرمول احتمال بینر می‌گوییم.

مثال ۳۱: احتمال افزایش قیمت سهام یک شرکت خصوصی در بورس در طول ۲ ماه مهر و آبان برابر 0.6 و احتمال سقوط قیمت سهام آن در طول این ۲ ماه برابر 0.3 است. همچنین اگر قیمت سهام شرکت در طول یکی از این دو ماه افزایش و در ماه دیگر کاهش پیدا کند در این صورت احتمال

اینکه در ماه اول قیمت سهام افزایش پیدا کند برابر $\frac{1}{4}$ می‌شود. مطلوبست احتمال اینکه قیمت سهام شرکت در ماه دوم هم افزایش پیدا کند به شرطی که در ماه اول افزایش پیدا کرده باشد.

حل: برای فضای نمونه ۴ حالت زیر را داریم:

$$S\{(I, I), (D, D), (I, D), (D, I)\}$$

I: افزایش قیمت سهام:

D: کاهش قیمت سهام:

توجه: برای سادگی حالت ثابت ماندن قیمت سهام در نظر نمی‌گیریم.

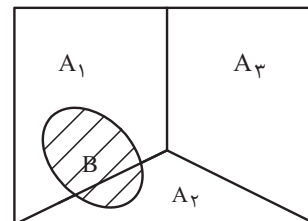
مثلاً (I-D) یعنی قیمت سهام در ماه اول افزایش و در ماه دوم کاهش یافته است. حال پیشامدهای زیر را تعریف می‌نیم:

$$A_1 = \{(I, I)\} = \text{پیشامد اینکه قیمت سهام در هر ماه افزایش پیدا می‌کند}$$

$$A_2 = \{(I, D), (D, I)\} = \text{پیشامد اینکه قیمت سهام در یک ماه افزایش و در یک ماه کاهش داشته باشد}$$

$$A_3 = \{(D, D)\} = \text{پیشامد اینکه قیمت سهام در هر دو ماه کاهش داشته باشد}$$

$$B = \text{پیشامد اینکه قیمت سهام در هر ماه اول افزایش داشته باشد}$$



$P(A_1|B) = P(\text{افزایش در ماه اول} | \text{افزایش در ماه دوم}) = P(\text{افزایش در ماه اول} | \text{افزایش در ماه دوم}) = P$. بنابراین می‌بایستی مقدار احتمال $P(A_1|B)$ را محاسبه کنیم با توجه به فرمول بینر داریم:

$$P(A_1|B) = \frac{P(A_1) P(B|A_1)}{P(A_1) P(B|A_1) + P(A_2) P(B|A_2) + P(A_3) P(B|A_3)}$$

$$P(A_1) = 0.6, \quad P(B|A_1) = 1, \quad P(B|A_2) = \frac{1}{2}, \quad P(B|A_3) = \frac{P(B \cap A_3)}{P(A_3)} = \frac{0}{P(A_3)} = 0$$

بنابراین داریم:

$$P(A_1|B) = \frac{0.6 \times 1}{0.6 \times 1 + 0.1 \times \frac{1}{2} + 0} = 0.923$$

فصل سوم متغیر های تصادفی و توابع توزیع

تعریف متغیر تصادفی

نتایج حاصل از یک آزمایش تصادفی را به صورت های مختلفی می توان بیان نمود . مثلاً شیر یا خط را با نماد 1 و -1 نیز نمایش داد . از آنجا که این اعداد در آزمایش به صورت تصادفی ظاهر می شوند ، می توان آنها را با یک متغیر تصادفی مثل x نمایش داد . با استفاده از متغیر تصادفی تفسیر نتایج حاصل از آزمایش های تصادفی ساده تر می شود .

تعریف :

یک متغیر تصادفی ، تابعی است از فضای نمونه به زیر مجموعه ای ناتمامی از اعداد حقیقی که آن را با حروف بزرگ مانند $X, Y, \dots$ نمایش می دهیم .
برای نمایش برد یک متغیر تصادفی یا مجموعه مقادیری که متغیر اختیار می کند ، از حروف کوچک مثل $x, y, \dots$ استفاده می کنیم .

مثال 1 :

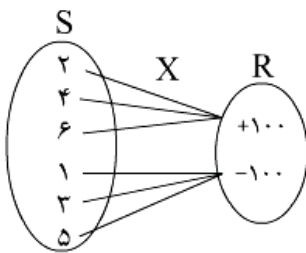
در یک بازی شخصی تاسی را می ریزد ، اگر عدد زوج بیاورد 100 تومان به او می دهیم و اگر عدد فرد بیاورد 100 تومان از او می گیریم . اگر متغیر تصادفی x در آمد شخص در نظر گرفته شود ، مقادیر آن را بدست آورید .

حل :

ابتدا فضای نمونه را بدست می آوریم .

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

همان طوری می بینید با توجه به نمودار زیر فضای نمونه S تحت تابع x به یک زیر مجموعه از R متشکل از دو عضو $+100$ و -100 برده شده است .



عدد $+100$ را برای دریافت و -100 را برای پرداخت پول در نظر می گیریم .

$$X: S \rightarrow R$$

$$\begin{array}{ll} X(2) = 100 & X(1) = -100 \\ X(4) = 100 & X(3) = -100 \\ X(6) = 100 & X(5) = -100 \end{array}$$

از آنجا که احتمال هر یک از اعداد 1 تا 6 برابر $\frac{1}{6}$ می باشد داریم :

$$A = \{ 2, 4, 6 \}$$

$$P(X = 100) = P(A) = \frac{n(A)}{n(S)} = \frac{3}{6} = \frac{1}{2}$$

$$B = \{ 1, 3, 5 \}$$

$$P(X = -100) = P(B) = \frac{n(B)}{n(S)} = \frac{3}{6} = \frac{1}{2}$$

حال می توانیم میانگین درآمد شخص را پس از تکرار بازی به دفعات محاسبه کنیم

$$\begin{array}{c|cc} X & 100 & -100 \\ \hline P & \frac{1}{2} & \frac{1}{2} \end{array} \Rightarrow \bar{X} = \frac{100 \times \frac{1}{2} + (-100) \times \frac{1}{2}}{\frac{1}{2} + \frac{1}{2}} = 0$$

با توجه به $\bar{X} = 0$ می توان نتیجه گرفت که شخص با تکرار بازی پولی به دست نمی آورد .



Jalase 1 – sco 2

انواع متغیر های تصادفی

متغیر های تصادفی عموماً بر 2 دسته می باشند :

1. گسسته : متغیر هایی که برد آن یا مجموعه مقادیری که اختیار می کند ، به صورت از هم جدا و یا شمارش پذیر از اعداد حقیقی باشند .
2. مقادیری اختیار می کند ، که در مجموعه ای پیوسته یا شمارش ناپذیر از اعداد حقیقی قرار دارند .

توجه :

همان طور که در فصل قبل اشاره کردیم احتمال وقوع هر عضو منفرد از مجموعه ای پیوسته برابر صفر می باشد . این مطلب برای متغیر های تصادفی پیوسته نیز برقرار می باشد . یعنی احتمال برابر بودن مقدار یک متغیر تصادفی پیوسته x بت هر عنصر منفرد از فضای برد آن برابر صفر می باشد .

مثال 2 :

در یک جعبه 22 مهره موجود می باشد که هر مهره دارای یکی از شماره های 1 و 2 و 3 و 4 و 5 و 6 می باشد . در ضمن به تعداد عدد نوشته شده بر روی هر مهره از همان مهره درون جعبه موجود می باشد . یک مهره از جعبه خارج می کنیم . اگر متغیر تصادفی x شماره مهره خارج شده باشد ، مطلوبست مقادیر احتمالی که متغیر x به خود می گیرد .

حل :

با توجه به صورت مثال یک مهره با شماره یک خواهیم داشت ، 2 مهره با شماره 2 و به همین ترتیب 6 مهره با شماره 6 داریم . فضای نمونه مقادیری که x اختیار می کند :

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

در نتیجه تابع احتمال به صورت زیر بدست می آید :

X	۱	۲	۳	۴	۵	۶
$P_X(x)$	$\frac{1}{21}$	$\frac{2}{21}$	$\frac{3}{21}$	$\frac{4}{21}$	$\frac{5}{21}$	$\frac{6}{21}$

مثلاً $P(X=6) = \frac{6}{21}$ است .

می توانیم به ازای هر مقدار x احتمال آن را به فرم $P_X(x)$ نشان داد که به آن تابع احتمال متغیر تصادفی x گوئیم .

تعریف : تابع احتمال متغیر تصادفی x ، تابعی است که از یک متغیر حقیقی x که با $P_X(x)$ نمایش داده می شود:

$$P_X(x) = P(X = x)$$

هر تابع احتمال برای متغیر تصادفی گسسته x باید در شرایط زیر صدق کند .

۱) $P_X(x) \geq 0$ ، به ازای هر x

۲) $\sum_{x \text{ برد}} P_X(x) = 1$



Jalase 1 – sco 3

توابع توزیع و چگالی برای متغیر تصادفی پیوسته

اگر $P_X(x)$ تابع احتمال متغیر تصادفی گسسته x باشد ، تابع توزیع $F_X(x)$ به صورت زیر تعریف می شود :

$$\forall x \in \mathbb{R} \quad ; \quad F_X(x) = p(X \leq x) = \sum_{t \leq x} P_X(t)$$

مثال 4 :

تابع توزیع را برای متغیر تصادفی که در مثال قبل حل شد را بدست آورید .

حل :

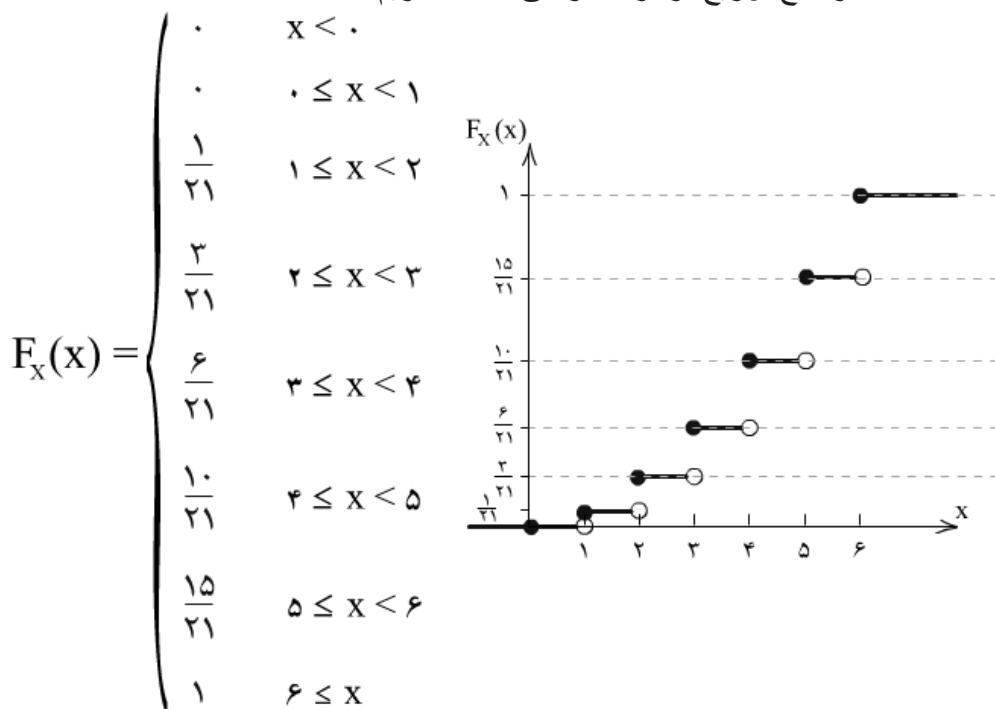
با توجه به جدول احتمالات متغیر تصادفی x داریم :

X	۱	۲	۳	۴	۵	۶
$P_X(x)$	$\frac{1}{21}$	$\frac{2}{21}$	$\frac{3}{21}$	$\frac{4}{21}$	$\frac{5}{21}$	$\frac{6}{21}$

$$F_X(1) = P(X \leq 1) = \frac{1}{21}$$

$$F_X(2/5) = P(X \leq 2/5) = \sum_{t \leq 2/5} P_X(t) = P_X(1) + P_X(2) = \frac{3}{21}$$

به ازای هر $x < 0$ مقدار تابع توزیع برابر صفر می باشد . داریم :



همان طور که مشاهده می شود تعریف تابع توزیع متغیر تصادفی x مشابه تعریف فراوانی تجمعی نسبی تعریف شده در فصل 1 است . مقدار فراوانی تجمعی قبل از 1 برابر 0 از 1 تا نزدیک 2 ،

$\frac{1}{21}$ الی آخر است .



Jalase 2 – sco 1

خواص توابع توزیع

توابع توزیع دارای خواص مشترکی می باشند . که با توجه به تعریف آنها بدست می آید . این خصوصیات در مسائل قبل نیز مشاهده می شوند که عبارتند از :

$$0 \leq F_X(x) \leq 1 \quad ; \quad \forall x \in \mathbb{R} \quad (1)$$

$$\lim_{x \rightarrow +\infty} F_X(x) = 1 \quad ; \quad \lim_{x \rightarrow -\infty} F_X(x) = 0 \quad (2)$$

$$F_X(\infty) = 1 \quad F_X(-\infty) = 0$$

(3) توابع توزیع همواره صعودی (غیر اکید) می باشد.

$$\forall a \leq b \rightarrow F_X(a) \leq F_X(b)$$

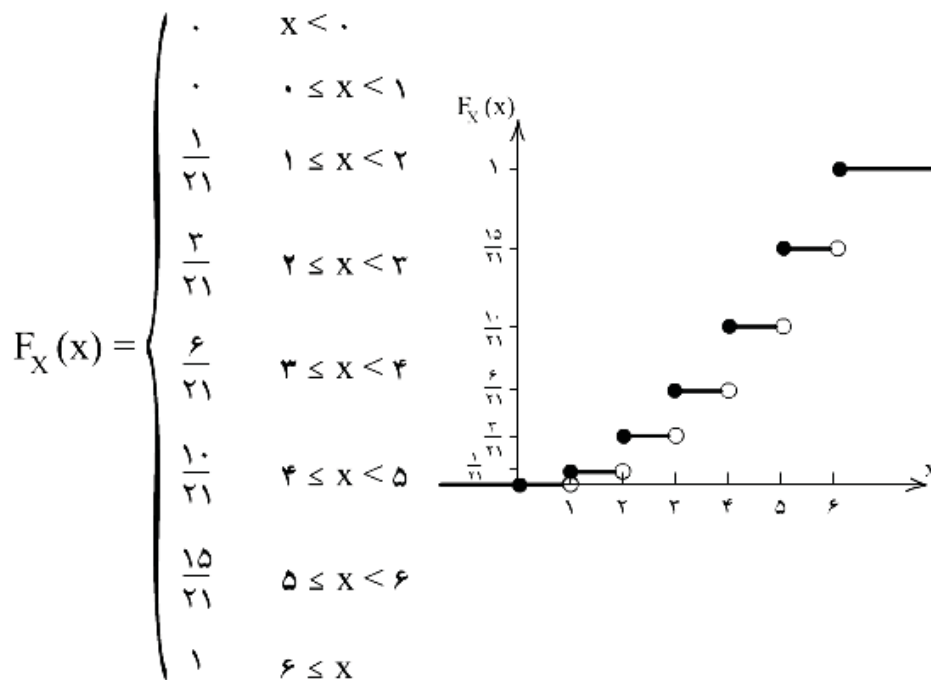
(4) تابع توزیع $F(x)$ در تمام نقاط حداقل از سمت راست، پیوسته می باشد.

$$\lim_{h \rightarrow 0^+} F_X(x+h) = F_X(x)$$



Jalase 2 – sco 2

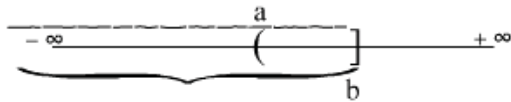
همواره از روی یک تابع توزیع متغیر تصادفی می توان مقادیر احتمال را محاسبه نمود.
به عنوان مثال اگر بخواهیم در مثال قبل مقدار $P(x=4)$ را محاسبه کنیم، با توجه به گسسته بودن متغیر تصادفی x داریم:



$$P(X=4) = \lim_{h \rightarrow 0} P(4-h < x \leq 4) \quad 1-3$$

در این صورت $P(x=4)$ را خواهیم داشت. برای محاسبه $P(a < x \leq b)$ به صورت زیر عمل می کنیم.

$$P(a < X \leq b) = P(X \leq b) - P(X \leq a) = F_X(b) - F_X(a)$$



توجه: برای سادگی در نوشتن می توان اندیس x را از تابع $F_X(x)$ حذف نمود و نوشت $F(x)$. البته اگر به شرطی که بدانیم با متغیر تصادفی x کار می کنیم. بنابراین:

$$P(a < X \leq b) = F(b) - F(a)$$

و داریم:

$$P(X = 4) = \lim_{h \rightarrow 0} P(4 - h < x \leq 4) \quad 1-3$$

$$\lim_{h \rightarrow 0} [F(4) - F(4 - h)] = F(4) - \lim_{h \rightarrow 0} F(4 - h)$$

که $\lim_{h \rightarrow 0} F(4 - h)$ برابر حد چپ تابع F_X در نقطه 4 می باشد و با $F(4^-)$ نمایش می دهیم. حال با توجه به ضابطه تابع F_X :

$$F(4^-) = \frac{6}{21}$$

و داریم:

$$P(X = 4) = F(4) - F(4^-) = \frac{10}{21} - \frac{6}{21} = \frac{4}{21}$$

درستی حاصل عبارت بالا از روی تابع احتمال $P_X(x)$ نیز آشکار است.



Jalase 2 – sco 3

در حالت کلی می توان رابطه زیر را نوشت:

$$P(X = b) = F(b) - F(b^-)$$

مقدار $P(x = b)$ برابر با میزان جهش نمودار تابع F_X در نقطه b می باشد. همچنین تمامی حالات دیگر را با توجه به فرمول بالا و تعریف تابع توزیع می توان بدست آورد.

$$۱) P(a < X \leq b) = P(X \leq b) - P(X \leq a) = F(b) - F(a)$$

$$۲) P(a \leq X < b) = P(X < b) - P(X < a) = F(b^-) - F(a^-)$$

$$۳) P(a < X < b) = P(X < b) - P(X \leq a) = F(b^-) - F(a)$$

$$۴) P(a \leq X \leq b) = P(X \leq b) - P(X < a) = F(b) - F(a^-)$$

$$۵) P(X > a) = ۱ - P(X \leq a) = ۱ - F(a)$$

$$۶) P(X < a) = F(a^-)$$

حد چپ تابع توزیع در نقطه a



Jalase 2 – sco 4

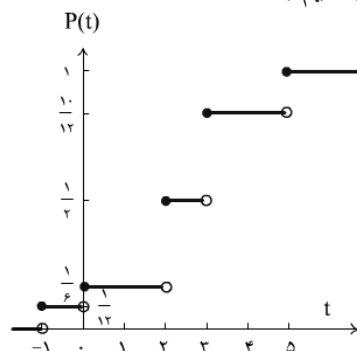
مثال 5 :

تابع توزیع برای متغیر تصادفی گسسته T به صورت زیر است . تابع احتمال آن را بدست بیاورید و نمودار آن را رسم کنید.

$$F(t) = \begin{cases} 0 & t < -۱ \\ \frac{1}{12} & -۱ \leq t < 0 \\ \frac{1}{6} & 0 \leq t < ۲ \\ \frac{1}{2} & ۲ \leq t < ۳ \\ \frac{10}{12} & ۳ \leq t < ۵ \\ 1 & ۵ \leq t \end{cases}$$

حل:

ابتدا نمودار تابع توزیع را رسم می کنیم .



مقادیری که تابع احتمال $P(T=t)$ می پذیرد ، در نقاط انفصال تابع $F(t)$ رخ می دهد . زیرا در نقاط دیگر مثل $T=1$ داریم :

$$P(T=1) = P(1 \leq T \leq 1) = F(1) - F(1^-) = \frac{1}{6} - \frac{1}{6} = 0$$

مقدار $P(t)$ را در نقطه انفصال تابع توزیع $F(t)$ به دست می آوریم :

$$P(T=-1) = F(-1) - F(-1^-) = \frac{1}{12} - 0 = \frac{1}{12}$$

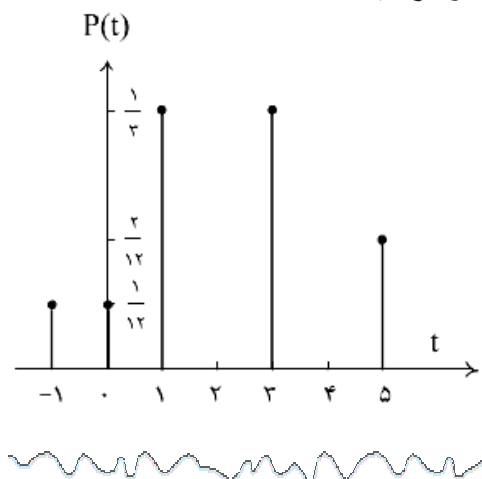
$$P(T=0) = F(0) - F(0^-) = \frac{1}{6} - \frac{1}{12} = \frac{1}{12}$$

$$P(T=2) = F(2) - F(2^-) = \frac{1}{2} - \frac{1}{6} = \frac{1}{3}$$

$$P(T=3) = F(3) - F(3^-) = \frac{10}{12} - \frac{1}{2} = \frac{4}{12} = \frac{1}{3}$$

$$P(T=5) = F(5) - F(5^-) = 1 - \frac{10}{12} = \frac{2}{12} = \frac{1}{6}$$

نمودار احتمال $P(t)$ شبیه به نمودار میله ای است .



Jalase 3 – sco 1

توابع چگالی

برای متغیر تصادفی گسسته X تابع چگالی که به فرم $f_X(x)$ نمایش داده می شود به صورت زیر تعریف می شود :

$$f_X(x) = P(X=x) \quad \text{عدد حقیقی } X$$

به عبارت دیگر تابع چگالی همان تابع احتمال متغیر تصادفی گسسته X می باشد که در رابطه زیر صدق کند :

$$F_X(x) = \sum_{\text{برد } X} f_X(x)$$

در برخی از کتاب ها به جای تابع احتمال برای متغیر های تصادفی گسسته مفهوم تابع چگالی را که دقیقاً همان تعریف را دارد به کار می برند .



Jalase 3 sco 2

توابع توزیع و چگالی برای متغیر تصادفی پیوسته

قبل از هر چیز به مثال زیر توجه کنید :

مثال :

فرض کنید مدت زمان توقف اتوبوس در ایستگاه و قبل از حرکت به صورت تصادفی 0 الی 15 دقیقه باشد . در این صورت متغیر تصادفی X را مدت زمان تأخیر حرکت اتوبوس پس از 5 دقیقه تعریف می کنیم . مطلوب است احتمال لین که اتوبوس به مدت 5 الی 6 دقیقه تأخیر داشته باشد .

حل :

در این مثال با یک فضای نمونه پیوسته مواجه ایم که می توان آن را به صورت زیر نمایش داد :

$$S = \{ 5 \leq x \leq 15 \}$$

نمایش تابعی متغیر تصادفی X عبارت است از $x_t = t - 5$ و $(t \in S)$ که میزان تأخیر

حرکت اتوبوس را نشان می دهد .

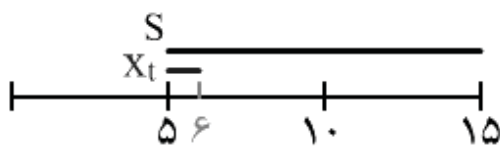
به این ترتیب می توان احتمال این که اتوبوس به مدت 5 الی 6 دقیقه تأخیر داشته باشد را محاسبه نمود .

احتمال این که $5 \leq x_t \leq 6$ باشد برابر است با طول بازه ی $[5, 6]$ تقسیم بر طول بازه

$[5, 15]$.

$$\frac{[5, 6]}{[5, 15]} = 0.1$$

لازم به تذکر است که در این جا تابع احتمال یک فاصله پیشامد مورد نظر به طول فاصله فضای نمونه ، تعریف شده است .



حال با توجه به این مثال به تعریف تابع چگالی ، تابع توزیع احتمال برای متغیر تصادفی پیوسته می پردازیم .

فرض کنید $x =$ متغیر تصادفی پیوسته باشد در این حالت نیز تعاریف تابع توزیع و چگالی در مورد متغیر x صدق می کند . با این تفاوت که در این حالت برای بدست آوردن تابع توزیع با توجه به پیوسته بودن از انتگرال گیری به جای جمع بستن روی تابع چگالی استفاده می کنیم . داریم :

$$F_X(t) = P(X \leq t) = \int_{-\infty}^t f_X(x) dx$$

که $f_X(x)$ تابع چگالی متغیر تصادفی پیوسته x است .
از این رابطه می توان تابع چگالی را بر حسب تابع توزیع بدست آورد ، به کمک مشتق گیری
یعنی

$$f_X(x) = \frac{d}{dx} F_X(x)$$

به عبارتی با مشتق گرفتن از تابع توزیع مقدار چگالی به دست می آید . با توجه به تعاریف می
توان خواص زیر را توابع توزیع و چگالی متغیر تصادفی پیوسته x بدست آورد.



Jalase 3 – sco 3

1- برای هر $x \in R$ ، داریم $f_X(x) \geq 0$.

توجه کنید که در حالت پیوسته $f_X(x)$ الزاما کوچکتر یا مساوی یک نمی باشد زیرا همان طور
که قبلا اشاره کردیم مقدار تابع احتمال برای متغیر تصادفی در یک نقطه صفر می باشد و
احتمالات می بایستی در این حالت در یک بازه محاسبه شوند .
-2

$$\int_{-\infty}^{+\infty} f_X(x) dx = 1$$

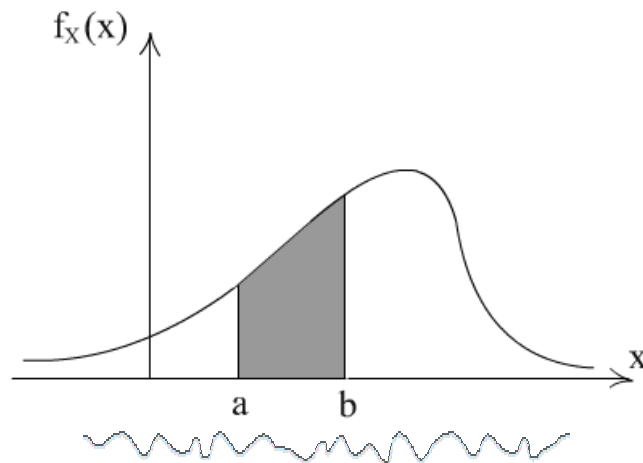
برای محاسبه احتمال قرار گرفتن متغیر تصادفی x در یک بازه به این ترتیب عمل می کنیم :

$$\begin{aligned} P(a < X \leq b) &= F(b) - F(a) = \int_{-\infty}^b f_X(x) dx - \int_{-\infty}^a f_X(x) dx \\ &= \int_a^b f_X(x) dx \end{aligned}$$

از طرفی در حالت پیوسته احتمال این که x برابر عدد دیگری باشد صفر می باشد . بنابراین می
توان نوشت:

$$\begin{aligned} P(a < X < b) &= P(a \leq X < b) = P(a < X \leq b) = P(a \leq X \leq b) \\ &= \int_a^b f_X(x) dx \end{aligned}$$

از لحاظ هندسی احتمال قرار گرفتن متغیر تصادفی پیوسته x در بازه a تا b برابر است با سطح
زیر منحنی تابع چگالی از a تا b . مطابق شکل زیر :



Jalase 4 – sco 1

مثال :

تابع چگالی احتمال متغیر تصادفی پیوسته x به صورت زیر می باشد :

$$f_X(x) = \begin{cases} x & 0 \leq x \leq 1 \\ c - x & 1 \leq x \leq 2 \\ 0 & \text{سایر مقادیر} \end{cases}$$

الف (مقدار متغیر c را بدست آورید .ب (تابع توزیع متغیر تصادفی x را بدست آورید .ج (احتمال $P(\frac{1}{2} \leq x \leq \frac{3}{2})$ را یکبار با استفاده از تابع چگالی $f_X(x)$ و یکبار با استفاده از تابعتوزیع $F_X(x)$ در نقطه x محاسبه کنید .

حل :

الف (برای بدست آوردن متغیر c می دانیم مساحت زیر منحنی f_X یعنی تابع چگالی می بایستی

برابر واحد باشد . بنابراین :

$$\int_{-\infty}^{\infty} f_X(x) dx = 1 \Rightarrow \int_{-\infty}^{\infty} f_X(x) dx = \int_{-\infty}^0 0 dx + \int_0^1 x dx + \int_1^2 (c - x) dx =$$

$$= \left. \frac{x^2}{2} \right|_0^1 + \left(cx - \frac{x^2}{2} \right) \Big|_1^2 = \frac{1}{2} + (2c - 2) - c + \frac{1}{2}$$

$$= \frac{1}{2} + 2c - 2 - c + \frac{1}{2} = c - 1$$

$$\Rightarrow c - 1 = 1 \Rightarrow c = 2$$

Jalase 4 – sco 2

ب (مقدار تابع توزیع f_X را محاسبه می کنیم . می دانیم :

$$F_X(x) = \int_{-\infty}^x f_X(t) dt$$

برای حل می بایستی این مقدار را در چهار حالت $x < 0$ و $0 \leq x < 1$ و $1 \leq x < 2$ و $x \geq 2$.

اگر $x < 0$ $\Rightarrow F_X(x) = \int_{-\infty}^x f_X(t) dt = 0$

اگر $0 \leq x < 1$ $\Rightarrow F_X(x) = \int_{-\infty}^x f_X(t) dt + \int_0^x t dt$

$$= 0 + \frac{t^2}{2} \Big|_0^x = \frac{x^2}{2}$$

اگر $1 \leq x < 2$ $\Rightarrow F_X(x) = \int_{-\infty}^x f_X(t) dt + \int_0^1 f_X(t) dt + \int_1^x f_X(t) dt$

$$= 0 + \int_0^1 t dt + \int_1^x (2-t) dt =$$

$$= 0 + \frac{t^2}{2} \Big|_0^1 + \left(2t - \frac{t^2}{2} \right) \Big|_1^x$$

$$= 0 + \frac{1}{2} + \left(2x - \frac{x^2}{2} \right) - \left(2 - \frac{1}{2} \right) =$$

$$= 0 + \frac{1}{2} + 2x - \frac{x^2}{2} - 2 + \frac{1}{2} = 2x - \frac{x^2}{2} - 1$$

اگر $2 \leq x$ $\Rightarrow F_X(x) = \int_{-\infty}^x f_X(t) dt + \int_0^1 f_X(t) dt + \int_1^2 f_X(t) dt + \int_2^x f_X(t) dt =$

$$\Rightarrow F_X(x) = 0 + \int_0^1 t dt + \int_1^2 (2-t) dt + \int_2^x f_X(t) dt$$

$$= 0 + \frac{t^2}{2} \Big|_0^1 + \left(2t - \frac{t^2}{2} \right) \Big|_1^2 + 0$$

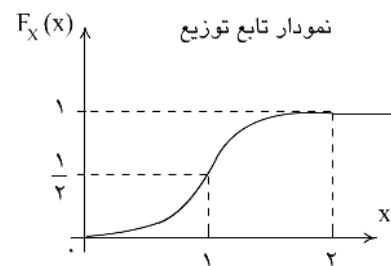
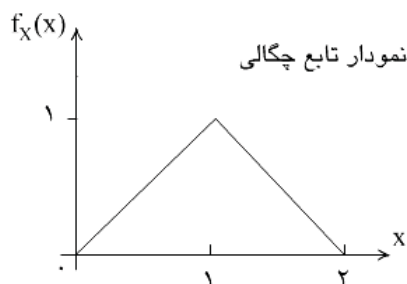
$$= \frac{1}{2} + (2 - 2) - \left(2 - \frac{1}{2} \right) = \frac{1}{2} + \frac{1}{2} = 1$$

بنابراین ضابطه تابع توزیع $F_X(x)$ برابر است با :

$$F_X(x) = \begin{cases} 0 & x < 0 \\ \frac{x^2}{2} & 0 \leq x < 1 \\ 2x - \frac{x^2}{2} - 1 & 1 \leq x < 2 \\ 1 & 2 \leq x \end{cases}$$

حال به نمودار دو تابع چگالی و توزیع توجه کنید :

نمودار تابع چگالی به صورت یک مثلث که قاعده پایینی آن بر روی محور x ها حداثصل 0 تا 2 است ، می باشد .



Jalase 4 – sco 3

ج (ابتدا مقدار $P(\frac{1}{2} \leq x \leq \frac{3}{2})$ را با استفاده از تابع چگالی $f_X(x)$ بدست می آوریم .

$$P(\frac{1}{2} \leq X \leq \frac{3}{2}) = \int_{\frac{1}{2}}^{\frac{3}{2}} f_X(t) dt = \int_{\frac{1}{2}}^1 t dt + \int_1^{\frac{3}{2}} (2 - t) dt$$

و این به خاطر این است که ضابطه تعریف تابع چگالی از 1 تا $\frac{3}{2}$ با $\frac{1}{2}$ تا 1 متفاوت است .

$$= \left. \frac{t^2}{2} \right|_{\frac{1}{2}}^1 + \left(2t - \frac{t^2}{2} \right) \Big|_1^{\frac{3}{2}} = \frac{1}{2} - \frac{1}{8} + \left(3 - \frac{9}{8} \right) - \left(2 - \frac{1}{2} \right)$$

$$= 2 - \frac{1}{8} = \frac{16}{8} - \frac{1}{8} = \frac{15}{8} = 0.75$$

مجدادا مقدار $P(\frac{1}{2} \leq x \leq \frac{3}{2})$ را با استفاده از تابع توزیع $F_X(x)$ که در بند ب محاسبه شده است بدست می آوریم.

$$P\left(\frac{1}{2} \leq X \leq \frac{3}{2}\right) = F_X\left(\frac{3}{2}\right) - F_X\left(\frac{1}{2}\right)$$

$$= 3 - \frac{9}{8} - 1 - \frac{1}{8} = 2 - \frac{10}{8} = \frac{3}{4}$$

تابع توزیع

ملاحظه ی کنید که مقدار احتمال از هر دو راه حل برابر است:

$$= 2 - \frac{10}{8} = \frac{6}{8} = \frac{3}{4} = 0.75$$

تابع چگالی



Jalase4 - sco 4

مثال :

متغیر تصادفی پیوسته x دارای تابع توزیع $F_X(x)$ با ضابطه زیر می باشد :

$$F_X(x) = \begin{cases} 0 & x < 0 \\ c - e^{-ax} & x \geq 0 \end{cases}, \quad a > 0$$

الف (مقدار متغیر c را بدست آورید ؟

ب (تعیین کنید متغیر a چه مفادیری می تواند بگیرد ، به شرطی که $F_X(x)$ همچنان تابع توزیع باقی بماند .

ج (تابع چگالی احتمال متغیر تصادفی x را بدست بیاورید .

د (مقادیر احتمال $P(x \leq 1)$ و $P(0 \leq x \leq 2)$ و $P(x \geq \ln 10)$ را بدست آورید .

ه (به ازای چه مقداری از متغیر y مقدار $P(x \leq y) = \frac{1}{2}$ می باشد ؟

حل :

می دانیم شرط این که تابعی مثل $F_X(x)$ تابع توزیع باشد برابر است با :

$$0 \leq F_X(x) \leq 1 \quad \forall x$$

$$\lim_{x \rightarrow -\infty} F_X(x) = 0, \quad \lim_{x \rightarrow +\infty} F_X(x) = 1$$

$$F_X(a) \leq F_X(b) \quad \forall a \leq b$$

$$\lim_{h \rightarrow 0} F_X(x+h) = F_X(x) \quad \forall x$$

پس از آن باقی شروط را بررسی می کنیم .

$$1 - \lim_{x \rightarrow +\infty} c - e^{-ax} = c - \lim_{x \rightarrow +\infty} e^{-ax} = c \rightarrow c = 1$$

اگر $x < 0$ آنگاه $F_X(x) = 0$ می باشد ، که در نتیجه شرط برقرار است .

اگر $x \geq 0$ آنگاه $F_X(x) = 1 - e^{-ax}$ داریم :

$$\begin{aligned} x \geq 0 \Rightarrow 0 \leq e^{-ax} \leq 1 &\rightarrow -1 \leq e^{-ax} < 0 \rightarrow 1 - 1 \leq e^{-ax} < 1 \\ &\rightarrow 0 \leq 1 - e^{-ax} < 1 \rightarrow 0 \leq F_X(x) \leq 1 \end{aligned}$$

$$\forall x \leq y \Rightarrow e^{-ax} \geq e^{-ay} \rightarrow -e^{-ax} \leq -e^{-ay}$$

$$\Rightarrow 1 - e^{-ax} \leq 1 - e^{-ay} \rightarrow F_X(x) \leq F_X(y)$$



Jalase 4 – sco 5

همچنین به ازای هر $x \geq 0$ تابع e^{-ax} پیوسته می باشد . بنابراین تابع $1 - e^{-ax}$ نیز حداقل از سمت راست پیوستگی دارد و در نقطه $x=0$ داریم :

$$\lim_{x \rightarrow 0^+} 1 - e^{-ax} = F_X(0) = 0$$

بنابراین تابع در نقطه 0 پیوسته می باشد . با توجه به برقراری هر 4 شرط می توانیم بگوییم $F_X(x)$ با ضابطه زیر یک تابع توزیع احتمال می باشد .

$$F_X(x) = \begin{cases} 0 & x < 0 \\ 1 - e^{-ax} & x \geq 0 \end{cases}$$

ب) از آنجایی که تابع $F_X(x)$ مستقل از مقدار a در شرایط تابع توزیع صدق می کند ، بنابراین به ازای همه مقادیر $a > 0$ تابع $F_X(x)$ نیز تابع توزیع می باشد .

ج) برای بدست آوردن تابع چگالی $f_X(x)$ کافی است از تابع توزیع مشتق بگیریم به این ترتیب که :

$$\text{اگر } x < 0 \quad F'_X(x) = 0$$

$$\text{اگر } x \geq 0 \quad F'_X(x) = 0 - (-a) e^{-ax} = a e^{-ax}$$

توجه کنید که تابع توزیع در نقطه $x=0$ پیوسته می باشد . اما تابع در این نقطه مشتق ندارد زیرا مشتق چپ و راست در $x=0$ با هم برابر نمی باشند .

$$\Rightarrow \text{مشتق چپ در صفر} \quad F'_X(0^-) = 0 \quad x < 0$$

$$\Rightarrow \text{مشتق راست در صفر} \quad F'_X(0^+) = a e^{-a(0)} = a \quad x \geq 0$$

بنابر این تابع چگالی احتمال برابر است با :

$$f_X(x) = \begin{cases} ae^{-ax} & x > 0 \\ 0 & \text{سایر مقادیر} \end{cases}$$

Jalase 4 – sco 6

(د)

$$P(x \leq 1) = F_X(1) = \int_{-\infty}^1 f_X(x) dx = 1 - e^{-a}$$

$$P(0 \leq x \leq 2) = F_X(2) - F_X(0) = 1 - e^{-2a} - (1 - 1) = 1 - e^{-2a}$$

$$\begin{aligned} P(x \geq \ln 10) &= 1 - P(x < \ln 10) = 1 - F_X(\ln 10) = \\ &= 1 - 1 + e^{-a \ln 10} = e^{-a \ln 10} = 10^{-a} = \frac{1}{10^a} \end{aligned}$$

ه (به فرم زیر عمل می کنیم :

$$P(X \leq y) = \frac{1}{10} \Rightarrow 1 - e^{-ay} = \frac{1}{10}$$

$$\Rightarrow \frac{1}{10} = e^{-ay} \rightarrow \ln \frac{1}{10} = -ay \rightarrow y = \frac{-1}{a} \ln \frac{1}{10}$$

فصل چهارم

امید ریاضی

۴-۱ با توجه به مطالب فصل اول می‌دانیم که اگر مقادیر $x_1 \dots x_n$ را از یک جامعه آماری داشته باشیم میانگین آنها را با $\bar{x} = \frac{\sum_{i=1}^n f_i x_i}{\sum_{i=1}^n f_i}$ نمایش

می‌دهیم حال فرض کنید مقادیر $x_1 \dots x_n$ مقادیری باشند که متغیر تصادفی گسسته X با مقدار احتمال $f(x_i)$ بخود می‌گیرد. در این صورت امید ریاضی متغیر تصادفی X را همان میانگین مقادیر $x_1 \dots x_n$ با فراوانی $f(x_i)$ در نظر می‌گیریم و با $E[X]$ یا μ_x به صورت زیر نمایش

می‌دهیم:

$$E[X] = \sum_{\text{برد } x} f(x_i) x_i$$

توجه کنید که در این حالت $\sum_{\text{برد } x} f_i = 1$ زیرا $f(x_i)$ خود تابع احتمال متغیر تصادفی X می‌باشد. با توجه به مطالب فوق می‌توان نتیجه گرفت که

امید ریاضی در واقع همان میانگین یا مقدار مورد انتظار متغیر تصادفی X است یعنی اگر تحت شرایط یکسان یک آزمایش را با توجه به مقادیر احتمال متغیر تصادفی X تکرار کنیم انتظار داریم چه مقداری از متغیر تصادفی X را مشاهده کنیم. برای روشن شدن مطلب به مثال زیر توجه کنید:

۴-۲ مثال ۸: یک شرکت بیمه انواع اتومبیل‌ها را بیمه بدنه می‌کند. فرض کنید در طول یک سال ۲۰ درصد افراد هیچگاه تصادف نکنند، ۳۰ درصد افراد تصادفاتی با مجموع هزینه ۱۰۰ هزار تومان برای بیمه، ۴۰ درصد افراد تصادفاتی با مجموع هزینه ۱ میلیون تومان برای بیمه و ۱۰ درصد افراد تصادفاتی با مجموع هزینه ۱۰ میلیون تومان برای بیمه داشته باشند. در این صورت انتظار داریم شرکت بیمه بطور متوسط در طول یک سال چه هزینه‌ای را برای بیمه بدنه اتومبیل‌ها پرداخت کند؟

حل: ابتدا متغیر تصادفی X را احتمال تصادف اتومبیل‌ها و مقادیر آنرا برابر با هزینه تقبل شده توسط شرکت بیمه در نظر می‌گیریم بنابر این داریم:

$$f(x) = \begin{cases} 0.2 & x = x_1 = 0 \\ 0.2 & x = x_2 = 100/000 \\ 0.4 & x = x_3 = 1/000/000 \\ 0.1 & x = x_4 = 10/000/000 \end{cases}$$

حال امید ریاضی متغیر تصادفی X را که برابر است با مقدار متوسط هزینه پرداخت شده توسط شرکت بیمه، محاسبه می‌کنیم:

$$E[X] = \sum_{i=1}^4 f(x_i) x_i = 0.2 \times 0 + (100/000) \cdot 0.2 + (1/000/000) \cdot 0.4 + (10/000/000) \cdot 0.1$$

$$= 30/000 + 400/000 + 1/000/000 = 1/430/000$$

به این ترتیب شرکت بیمه می‌بایستی سالانه مبلغ ۱/۴۳۰/۰۰۰ تومان را به ازای بیمه بدنه اتومبیل‌ها پرداخت کند.

۴-۳ مثال ۹: فرض کنید در مثال قبل شرکت بیمه تعداد ۱۰۰ اتومبیل را بیمه بدنه کرده باشد در این صورت حداقل قیمت پیشنهادی شرکت برای هر اتومبیل چقدر باشد تا شرکت ضرر نکند؟

حل: توجه کنید که با توجه به مقدار امید ریاضی متغیر تصادفی X که در مثال قبل درست آمد می‌دانیم شرکت سالانه ۱/۴۳۰/۰۰۰ تومان هزینه می‌کند. با تقسیم این مبلغ بر تعداد اتومبیل‌ها میزان هزینه به ازای هر اتومبیل بدست می‌آید که برابر است با:

$$\frac{1/430/000}{100} = 14/300 \text{ میزان هزینه سالانه هر اتومبیل}$$

حال برای آنکه شرکت ضرر ندهد می‌بایستی به ازای بیمه بدنه هر اتومبیل حداقل مبلغ ۱۴/۳۰۰ تومان را دریافت کند.

۴-۴ برخی از خواص امیر ریاضی

از آنجا که امید ریاضی بر پایه میانگین $\bar{X}$ تعریف شده است خواص بدست آمده برای میانگین در مورد امید ریاضی نیز برقرار است که عبارتند از:

$$۱- E[ax] = aE[X] \quad \text{مقدار ثابت } A$$

$$۲- E[x+b] = E[X] + b \quad \text{مقدار ثابت } B$$

$$E[ax+b] = aE[X] + b \quad \text{و در حالت کلی:}$$

۴-۵ امید ریاضی تابعی از یک متغیر تصادفی

تا بحال امید ریاضی را بر حسب متغیر تصادفی X محاسبه کردیم اما می‌توان آنرا بر حسب تابعی مثل X^2 ، $X-2$ ، محاسبه نمود به این ترتیب می‌توان طیف وسیعتری از مسائل را حل نمود. برای محاسبه امید ریاضی بر حسب تابعی مثل $g(X)$ از رابطه زیر استفاده می‌کنیم:

$$E[g(x)] = \sum_{\text{برد } x} g(x) f_x(x)$$

به عنوان مثال $E[X]$ در واقع همان $E[g(x)=x]$ می‌باشد.

واریانس و انحراف معیار یک متغیر تصادفی X نیز با قرار دادن $g(X) = (X-\mu)^2$ بدست می‌آید و داریم:

$$\delta_X^2 = E\left[(X-\mu)^2\right]$$

و انحراف معیار متغیر تصادفی X با جذر از δ_X^2 بدست می‌آید $\delta_X = \sqrt{\delta_X^2}$ می‌توان نشان داد که δ_X^2 از رابطه زیر بدست می‌آید:

$$\begin{aligned} \delta_X^2 &= E[X^2] - E^2[X] \\ \delta_X^2 &= E[(X-E[X])^2] = E[X^2 - 2XE[X] + E^2[X]] \quad \text{اثبات:} \\ &= E[X^2] - E[2XE[X]] + E^2[X] \\ &= E[X^2] - E[2E[X]]E[X] + E^2[X] \\ &= E[X^2] - E^2[X] \end{aligned}$$

۴-۶ گشتاورها

اگر امید ریاضی را بر حسب تابع $g(X) = (X-a)^k$ محاسبه کنیم به آن گشتاور مرتبه k ام حول a گویند به عبارتی:

$$a \text{ حول } = E[(X-a)^k] = \text{گشتاور مرتبه } k \text{ام}$$

همچنین اگر $a = \bar{X}$ به آن گشتاور مرتبه k ام مرکزی گویند و اگر $a = 0$ باشد به آن گشتاور مرتبه k ام حول صفر گویند. که آنرا با m_k نمایش می‌دهند.

$$\text{گشتاور مرتبه } k \text{ام مرکزی} = E[(X-\bar{X})^k]$$

$$m_k = \text{گشتاور مرتبه } k \text{ام حول صفر} = E[X^k]$$

۴-۷ مثال ۱۰: مقادیر m_1 و m_2 را بر حسب μ و δ_X^2 بدست بیاورید:

$$m_1 = E[X^1] = \mu_X$$

$$m_2 = E[X^2] = \delta_X^2 + \mu_X^2$$

حل:

بنابراین با داشتن مقادیر گشتاورهای مرتبه اول و دوم می توان مقادیر امید ریاضی و واریانس را بدست آورد. همچنین با داشتن مقادیر سایر گشتاورهای مرتبه سوم به بعد نیز می توان اطلاعاتی در مورد متغیر تصادفی X بدست آورد.

۴-۸ تابع مولد گشتاورها

تابع مولد گشتاورها برای متغیر تصادفی X بصورت زیر تعریف می شود و با $m_X(t)$ نمایش داده می شود:

$$m_X(t) = E[e^{tX}]$$

$m_X(t)$ را به این دلیل تابع مولد گشتاورها می نامیم که با k بار مشتق گیری از آن نسبت به t و قرار دادن $t=0$ مقدار گشتاور مرتبه k ام متغیر تصادفی X بدست می آید.

به عنوان مثال با یکبار مشتق گیری نسبت به t داریم: (با فرض اینکه بتوان ترتیب اعمال مشتق گیری و امیدگیری را جابجا نمود)

$$m_X^{(1)}(t) = \frac{d}{dt} E[e^{tX}] = E\left[\frac{d}{dt} e^{tX}\right] = E[X e^{tX}]$$

$$m_X^{(1)}(0) = E[X e^{0X}] = E[X] = m_1$$

با قرار دادن $t=0$ داریم:

که همان گشتاور مرتبه اول یا μ_X می باشد. به همین ترتیب با K بار مشتق گیری داریم:

$$m_X^{(k)}(t) = E\left[\frac{d^k}{dt^k} e^{tX}\right] = E[X^k e^{tX}]$$

$$m_X^{(k)}(0) = E[X^k]$$

و در $t=0$ خواهیم داشت:

که همان گشتاور مرتبه k ام متغیر X می باشد.

۴-۹ امید ریاضی برای متغیر تصادفی پیوسته

برای متغیر تصادفی پیوسته X امید ریاضی به جای جمع بستن روی مقادیر $x f(x)$ با انتگرال گیری روی این مقادیر بدست می آید بصورت زیر:

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

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

همینطور داریم:

به همین ترتیب برای محاسبه گشتاورها و توابع مولد برای متغیر تصادفی پیوسته X از رابطه فوق در فرمول استفاده می شود.

۴-۱۰ مثال: فرض کنید برای متغیر تصادفی X مقدار تابع مولد گشتاور برابر با $m_X(t) = \frac{e^t - 1}{t}$ باشد در این صورت مقدار گشتاور مرتبه k ام

حول مبدأ یا $E[X^K]$ را محاسبه کنید.

حل: می دانیم مشتق مرتبه K ام تابع مولد گشتاور نسبت به t در نقطه صفر برابر با $E[X^K]$ می باشد بنابراین:

$$E[X^K] = \frac{d^k}{dt^k} \left(\frac{e^t - 1}{t} \right) \Big|_{t=0} \Rightarrow \frac{d}{dt} \left(\frac{e^t - 1}{t} \right) = \frac{te^t - e^t + 1}{t^2}$$

ملاحظه می‌کنید که مشتق مرتبه اول در نقطه $t=0$ موجود نمی‌باشد به همین ترتیب سایر مشتقات مراتب بالاتر نیز در نقطه $t=0$ موجود نمی‌باشند اما این به معنی عدم وجود $E[X^K]$ نمی‌باشد.

در مواقعی که مشتق در نقطه $t=0$ موجود نباشد می‌توان از قضیه زیر برای بدست آوردن گشتاور k ام حول مبدأ استفاده کرد.

قضیه: گشتاور مرتبه k ام متغیر تصادفی X حول مبدأ برابر است با ضریب $\frac{t^k}{k!}$ در بسط تیلور تابع مولد گشتاور متغیر تصادفی X :

$$m_X(t) = \sum_{k=0}^{\infty} E[X^K] \frac{t^k}{k!}$$

حال با استفاده از قضیه فوق مجدداً مقدار $E[X^K]$ را در مثال قبل محاسبه می‌کنیم. می‌دانیم: $e^t = \sum_{k=0}^{\infty} \frac{t^k}{k!}$ بنابراین:

$$\begin{aligned} m_X(t) &= \frac{e^t - 1}{t} = \frac{\sum_{k=0}^{\infty} \frac{t^k}{k!} - 1}{t} = \sum_{k=0}^{\infty} \frac{t^{k-1}}{k!} = \sum_{k=0}^{\infty} \frac{t^k}{(k+1)!} \\ &= \sum_{k=0}^{\infty} \underbrace{\frac{1}{k+1}}_{E[X^K]} \frac{t^k}{k!} \end{aligned}$$

بنابراین ملاحظه می‌کنید که ضریب $\frac{t^k}{k!}$ در بسط فوق برابر $\frac{1}{k+1}$ می‌باشد یعنی: $E[X^K] = \frac{1}{k+1}$.

۱۱-۴ مثال: تابع مولد گشتاور متغیر تصادفی X برابر است با $m_X(t) = \frac{1}{1-t}$ ، μ_X و δ_X^2 و $E[X^K]$.

حل: در این مثال مقدار مشتق مرتبه k ام $\frac{1}{1-t}$ در نقطه $t=0$ موجود می‌باشد اما محاسبه آن مشکل می‌باشد بنابراین برای راحتی از بسط تیلور

استفاده می‌کنیم:

$$m_X(t) = \frac{1}{1-t} = 1 + t + t^2 + t^3 + \dots = \sum_{k=0}^{\infty} t^k = \sum_{k=0}^{\infty} k! \frac{t^k}{k!}$$

بنابراین $E[X^K] = k!$.

حال مقادیر μ_X و δ_X^2 برابرند با:

$$\begin{aligned} \mu_X &= E[X] = 1! = 1 \\ \delta_X^2 &= E[X^2] - E[X]^2 = 2! - 1 = 1 \end{aligned}$$

۱۲-۴ مثال: متغیر تصادفی X دارای تابع چگالی زیر می‌باشد:

$$f_X(x) = \begin{cases} \frac{1}{2} e^{-\frac{x}{2}} & x > 0 \\ 0 & x \leq 0 \end{cases}$$

مطلوبست :

الف) محاسبه تابع توزیع متغیر تصادفی X .

ب) محاسبه μ_X ، δ_X^2 و گشتاور مرتبه k ام حول مبدأ و حول میانگین

حل: الف:

$$F_X(x) = \int_{-\infty}^x f_X(t) dt = \int_0^x \frac{1}{\tau} e^{-\frac{t}{\tau}} dt = -e^{-\frac{t}{\tau}} \Big|_0^x = -e^{-\frac{x}{\tau}} + 1$$

$$F_X(x) = \begin{cases} 1 - e^{-\frac{x}{\tau}} & x \geq 0 \\ 0 & x < 0 \end{cases}$$

بنابراین:

ب) ابتدا تابع مولد گشتاور را محاسبه می‌کنیم:

$$m_X(t) = E[e^{tx}] = \int_{-\infty}^{\infty} e^{tx} \frac{1}{\tau} e^{-\frac{x}{\tau}} dx = \int_0^{\infty} \frac{1}{\tau} e^{x(t-\frac{1}{\tau})} dx$$

$$= \frac{1}{\tau(t-\frac{1}{\tau})} e^{x(t-\frac{1}{\tau})} \Big|_0^{\infty} = \frac{1}{\tau(t-\frac{1}{\tau})} = \frac{1}{1-\tau t}$$

به شباهت بین تابع مولد گشتاور در ایم مثال و مثال قبل توجه کنید.

$$\mu_X = E[X] = \frac{d}{dt} \left(\frac{1}{1-\tau t} \right) \Big|_{t=0} = \tau$$

$$\delta_X^2 = E[X^2] - E^2[X] = \frac{d^2}{dt^2} \left(\frac{1}{1-\tau t} \right) \Big|_{t=0} = \tau^2$$

برای محاسبه گشتاور مرتبه k ام حول میانگین داریم:

$$E[(X-\mu_X)^K] = \text{گشتاور مرتبه } k \text{ام حول میانگین}$$

برای محاسبه $E[(X-\mu_X)^K]$ از تابع مولد گشتاور حول میانگین که به صورت زیر تعریف می‌شود استفاده می‌کنیم:

$$m_{X-\mu_X}(t) = E[e^{t(X-\mu_X)}]$$

حال توجه کنید که:

$$m_{X-\mu_X}(t) = E\left[e^{tX-t\mu_X}\right] = E\left[e^{tX} e^{-t\mu_X}\right] = e^{-t\mu_X} E\left[e^{tX}\right] = e^{-t\mu_X} m_X(t)$$

بنابراین رابطه کلی زیر بین گشتاور مرتبه k ام حول مبدأ و میانگین موجود است:

$$m_{X-\mu_X}(t) = e^{-t\mu_X} m_X(t)$$

حال به راحتی می‌توان $m_{X-\mu_X}(t)$ را با توجه به مقدار $m_X(t)$ بدست آورد داریم:

$$m_{X-\mu_X}(t) = e^{-\tau t} \frac{1}{1-\tau t}$$

بنابراین مقدار گشتاور مرتبه k ام حول میانگین برابر است با مشتق مرتبه k ام $m_{X-\mu_X}(t)$ نسبت به t در نقطه $t = 0$.

مقدار $E[X^K]$ برابر است با:

$$m_X(t) = \frac{1}{1-\gamma t} = 1 + \gamma t + (\gamma t)^2 + (\gamma t)^3 = \sum_{k=0}^{\infty} (\gamma t)^k$$

$$= \sum_{k=0}^{\infty} \gamma^k \frac{t^k}{k!} \Rightarrow E[X^K] = \gamma^K K!$$

۱۳-۴ مثال: تابع توزیع متغیر تصادفی X بصورت زیر داده شده است:

$$f(x) = \begin{cases} 0 & x < 0 \\ \frac{1}{2}x^2 & 0 \leq x < 1 \\ -1 + 2x - \frac{1}{2}x^2 & 1 \leq x < 2 \\ 1 & x \geq 2 \end{cases}$$

مطلوبست:

الف) تابع چگالی متغیر تصادفی X .

ب) μ_X و δ_X^2

حل: الف)

$$f(x) = \frac{d}{dx} f(x) \Rightarrow \text{اگر } x < 0 \quad f(x) = 0$$

$$\text{اگر } 0 \leq x < 1 \Rightarrow f(x) = 2 - x$$

$$x \geq 2 \Rightarrow f(x) = 0$$

بنابراین:

$$f(x) = \begin{cases} x & 0 \leq x < 1 \\ 2 - x & 1 \leq x < 2 \\ 0 & \text{سایر مقادیر} \end{cases}$$

ب)

$$E[X] = \int_{-\infty}^{\infty} f(x) dx = \int_0^1 x^2 dx + \int_1^2 x(2-x) dx =$$

$$\frac{x^3}{3} \Big|_0^1 + \left(x^2 - \frac{x^3}{3} \right) \Big|_1^2 = \frac{1}{3} + \left(4 - \frac{8}{3} \right) - \left(1 - \frac{1}{3} \right) = 1$$

$$E[X^2] = \int_{-\infty}^{\infty} x^2 f(x) dx = \int_0^1 x^3 dx + \int_1^2 x^2 (2-x) dx =$$

$$\frac{x^4}{4} \Big|_0^1 + \left(\frac{2x^3}{3} - \frac{x^4}{4} \right) \Big|_1^2 = \frac{1}{4} + \left(\frac{16}{3} - \frac{16}{4} \right) - \left(\frac{2}{3} - \frac{1}{4} \right) = \frac{7}{6}$$

$$\delta_X^2 = E[X^2] - E[X]^2 = \frac{7}{6} - 1 = \frac{1}{6}$$

فصل پنجم

۵-۱ متغیرهای تصادفی دو و چند بعدی

فرض کنید دو متغیر تصادفی X و Y را داشته باشیم می‌دانیم هر کدام از دو متغیر به هر عضو فضای نمونه S مقداری حقیقی و منحصر به فرد نسبت می‌دهند همینطور به ازای هر کدام از X و Y یک تابع چگالی احتمال $f_X(x)$ و $f_Y(y)$ خواهیم داشت.

اگر بخواهیم احتمال وقوع همزمان مقادیر برد هر یک از متغیرهای X و Y را بصورت یک تابع احتمال نشان دهیم، متغیر تصادفی $Z=(X, Y)$ را متغیر تصادفی دو بعدی با تابع احتمال $f_{X,Y}(x, y)$ معرفی می‌کنیم.

متغیر تصادفی دو بعدی تمامی خواص متغیرهای یک بعدی را دارا می‌باشد همینطور خواص زیر برای تابع چگالی احتمال آن برقرار است:

$$f_{X,Y}(x, y) = P(X=x, Y=y)$$

برای متغیر تصادفی دو بعدی گسسته:

$$1- \quad 0 \leq f_{X,Y}(x, y) \leq 1 \quad \forall x, y$$

$$2- \quad \sum_{Y \text{ برد}} \sum_{X \text{ برد}} f_{X,Y}(x, y) = 1$$

برای متغیر تصادفی دو بعدی پیوسته:

$$1- \quad f_{X,Y}(x, y) \geq 0 \quad \forall x, y$$

$$2- \quad \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} f_{X,Y}(x, y) = 1$$

$$3- \quad P(a \leq x \leq b, c \leq y \leq d) = \int_c^d \int_a^b f_{X,Y}(x, y) dx dy$$

برای روشن شدن مطلب به مثال زیر توجه کنید:

مثال ۱: یک عدد تاس را که بر روی سه وجه آن عدد ۱ و بر روی سه وجه دیگر عدد ۲ حک شده است را دو بار پرتاب می‌کنیم و متغیرهای X و Y را بصورت زیر تعریف می‌کنیم:

$$X = \text{مجموع دو عدد ظاهر شده} = \{2, 3, 4\}$$

$$Y = \text{تفاضل دو عدد ظاهر شده} = \{-1, 0, 1\}$$

بنابراین متغیر تصادفی دو بعدی (X, Y) را می‌توان به این ترتیب تعریف نمود:

زوج مرتب نمایش دهنده مجموع و تفاضل دو عدد ظاهر $(X, Y) =$ شده در دو بار پرتاب تاس.

چند عضو فضای نمونه S عبارتند از:

$$= \{(2, -1), (2, 0), (2, 1), (3, -1), \dots\}$$

به این ترتیب تابع احتمال دو بعدی $f_{X,Y}(x, y)$ بصورت زیر می‌باشد:

۵-۳

$Y \backslash X$	۲	۳	۴
-۱	۰	$\frac{1}{4}$	۰
۰	$\frac{1}{4}$	۰	$\frac{1}{4}$
۱	۰	$\frac{1}{4}$	۰

به عنوان مثال مقادیر احتمال بصورت زیر محاسبه می‌شود:

$$f_{X,Y}(x=2, y=-1) = 0$$

برای اینکه مجموع دو عدد ظاهر شده برابر با عدد ۲ باشد باید بر روی تاس در پرتاب اول عدد ۱ و در پرتاب دوم عدد ۱ ظاهر شده باشد. بنابراین احتمال این حالت برابر صفر می‌باشد.

اما برای حالت $(2, 0)$ این احتمال برابر است با:

$$f_{X,Y}(x=2, y=0) = \frac{1}{2} \times \frac{1}{2} = \frac{1}{4}$$

برای آنکه مجموع ۲ باشد می‌بایستی در پرتاب اول عدد ۱ و در پرتاب دوم عدد ۱ ظاهر شده باشد بنابراین احتمال اینکه در پرتاب اول عدد ظاهر شود

برابر $\frac{1}{2}$ می‌باشد و چون ظاهر شدن عدد ۱ در پرتاب دوم مستقل از پرتاب اول است بنابراین احتمال ظاهر شدن ۱ در هر پرتاب برابر است با:

$$\frac{1}{2} \times \frac{1}{2} = \frac{1}{4}$$

از روی تابع چگالی احتمال دو بعدی در این مثال می‌توان مقدار تابع چگالی احتمال متغیر تصادفی X را مستقل از Y محاسبه کنیم که آنرا تابع احتمال حاشیه‌ای یا کناری X می‌نامیم. که جمع بستن روی مقادیر Y در هر سطر بدست می‌آید به عبارتی:

$f_X(x) \sum_y f_{X,Y}(x,y)$	$Y \backslash X$	۲	۳	۴	$f_Y(y)$
	-۱	۰	$\frac{1}{4}$	۰	$\frac{1}{4}$
	۰	$\frac{1}{4}$	۰	$\frac{1}{4}$	$\frac{1}{2}$
	۱	۰	$\frac{1}{4}$	۰	$\frac{1}{4}$
$f_X(x)$		$\frac{1}{4}$	$\frac{1}{2}$	$\frac{1}{4}$	

به همین ترتیب تابع چگالی احتمال Y نیز با جمع بستن روی مقادیر X در جدول بدست می‌آید که به آن تابع احتمال حاشیه‌ای برای Y گویند.

$$f_Y(y) \sum_x f_{X,Y}(x,y)$$

۴-۵ برای نمونه $f_X(x)$ بصورت زیر بدست می‌آید:

$$f_X(2) = f(2, -1) + f(2, 0) + f(2, 1) = 0 + \frac{1}{4} + 0 = \frac{1}{4}$$

$$f_X(3) = f(3, -1) + f(3, 0) + f(3, 1) = \frac{1}{4} + 0 + \frac{1}{4} = \frac{1}{2}$$

$$f_X(4) = f(4, -1) + f(4, 0) + f(4, 1) = 0 + \frac{1}{4} + 0 = \frac{1}{4}$$

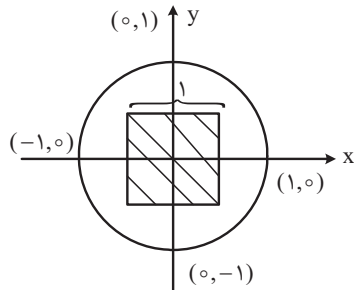
بنابراین:

$$f_X(x) = \begin{cases} \frac{1}{4} & x=2 \\ \frac{1}{2} & x=3 \\ \frac{1}{4} & x=4 \\ 0 & \text{سایر مقادیر} \end{cases}$$

صحت مقادیر $f_X(x)$ را با استفاده از تعریف متغیر تصادفی X نیز می‌توان بررسی نمود مثلاً $f_X(3)$ یعنی عدد ظاهر شده در پرتاب اول تاس ۱ و در پرتاب دوم ۲ بوده است یا عدد ظاهر شده در پرتاب اول ۲ و در پرتاب دوم ۱ بوده است که احتمال آن برابر است با:

$$f_X(2) = \frac{1}{2} \times \frac{1}{2} + \frac{1}{2} \times \frac{1}{2} = \frac{1}{2}$$

۵-۵ مثال ۲: از درون دایره واحد یک نقطه به تصادف انتخاب می‌کنیم متغیر تصادفی دو بعدی (X, Y) را طول و عرض نقطه انتخاب شده در نظر می‌گیریم مطلوبست محاسبه تابع چگالی احتمال دو بعدی $f_{X,Y}(x, y)$ و با استفاده از آن احتمال اینکه نقطه انتخاب شده درون مربع واحد قرار بگیرد را محاسبه کنید. همچنین تابع چگالی احتمال حاشیه‌ای برای X و Y را محاسبه کنید.



حل: با توجه به شکل روبرو مقدار تابع چگالی را بصورت زیر در نظر می‌گیریم:

$$f_{X,Y}(x, y) = \begin{cases} c & x^2 + y^2 \leq 1 \\ 0 & \text{سایر مقادیر} \end{cases}$$

حال مقدار مجهول C را محاسبه می‌کنیم:

$$\int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} f_{X,Y}(x, y) dx dy = 1$$

$$\Rightarrow \int_0^1 \int_0^{2\pi} c r d\theta dr \Rightarrow c\pi = 1 \rightarrow c = \frac{1}{\pi}$$

بنابراین تابع چگالی احتمال دو متغیر $f_{X,Y}(x, y)$ برابر است با:

$$f_{X,Y}(x, y) = \begin{cases} \frac{1}{\pi} & x^2 + y^2 \leq 1 \\ 0 & \text{سایر مقادیر} \end{cases}$$

برای آنکه نقطه انتخابی درون مربع واحد باشد باید داشته باشیم:

$$-\frac{1}{2} \leq x \leq \frac{1}{2}, \quad -\frac{1}{2} \leq y \leq \frac{1}{2}$$

مقدار احتمال با انتگرال‌گیری روی $f_{X,Y}(x, y)$ در بازه فوق بدست می‌آید:

$$p\left(-\frac{1}{2} \leq x \leq \frac{1}{2}, \frac{1}{2} \leq y \leq \frac{1}{2}\right) = \int_{-\frac{1}{2}}^{\frac{1}{2}} \int_{-\frac{1}{2}}^{\frac{1}{2}} \frac{1}{\pi} dx dy = \frac{1}{\pi}$$

توجه کنید که مقدار $\frac{1}{\pi}$ برابر با مساحت مربع تقسیم بر مساحت کل دایره می‌باشد که در فصل دوم نیز به این روش محاسبه می‌شد.

۵-۶ برای محاسبه تابع چگالی احتمال حاشیه‌ای برای متغیر X در حالت پیوسته به جای جمع بستن روی مقادیر Y از انتگرال‌گیری استفاده می‌کنیم به عبارتی:

$$\begin{aligned} f_X(x, y) &= \int_{-\infty}^{+\infty} f_{X,Y}(x, y) dy = \int_{-\sqrt{1-x^2}}^{\sqrt{1-x^2}} \frac{1}{\pi}(y) \left| \sqrt{1-x^2} \right. \\ &= \frac{1}{\pi} (\sqrt{1-x^2} - (-\sqrt{1-x^2})) = \frac{2\sqrt{1-x^2}}{\pi} \end{aligned}$$

به همین ترتیب مقدار $f_Y(y)$ نیز بدست می‌آید:

$$f_Y(y) = \int_{-\infty}^{+\infty} f_{X,Y}(x,y) dx = \int_{-\sqrt{1-y^2}}^{\sqrt{1-y^2}} \frac{1}{\pi}(x) \frac{\sqrt{1-y^2}}{-\sqrt{1-y^2}}$$

$$= \frac{1}{\pi} (\sqrt{1-y^2} - (-\sqrt{1-y^2})) = \frac{2\sqrt{1-y^2}}{\pi}$$

بنابراین خواهیم داشت:

$$f_X(x) = \begin{cases} \frac{2\sqrt{1-x^2}}{\pi} & -1 \leq x \leq 1 \\ 0 & \text{سایر مقادیر} \end{cases}$$

$$f_Y(y) = \begin{cases} \frac{2\sqrt{1-y^2}}{\pi} & -1 \leq y \leq 1 \\ 0 & \text{سایر مقادیر} \end{cases}$$

۵-۱.۴ تابع توزیع دو متغیره

تابع توزیع دو متغیره نیز کاملاً مشابه حالت یک متغیره بدست می‌آید به این ترتیب که:

$$F_{X,Y}(t_1, t_2) = p(x \leq t_1, y \leq t_2) \quad \forall (t_1, t_2) \in \mathbb{R}^2$$

حال اگر X و Y گسسته باشند تابع توزیع با جمع بستن روی X و Y تا (t_1, t_2) بدست می‌آید و اگر X و Y پیوسته باشند با انتگرال‌گیری تا (t_1, t_2) .

حل: با توجه به تعریف داریم:

$Y \backslash X$	۲	۳	۴
-۱	۰	$\frac{1}{4}$	$\frac{1}{4}$
۰	$\frac{1}{4}$	$\frac{1}{2}$	$\frac{3}{4}$
۱	$\frac{1}{4}$	$\frac{1}{4}$	۱

با توجه به تابع توزیع می‌توان مقدار

احتمال اینکه مجموع دو عدد ظاهر شده کمتر از ۴ و تفاضل دو عدد ظاهر شده کمتر از ۰ باشد را بدست آورد که برابر است با $f_{X,Y}(4,0)$:

$$p(x \leq 4, y \leq 0) = F_{X,Y}(4,0) = \frac{3}{4}$$

۵-۱.۴ امید ریاضی و گشتاورها برای توابع چگالی دو متغیره

همانند توابع چگالی یک متغیره امید ریاضی و توابع مولد گشتاورهای توام X و Y به صورت زیر محاسبه می‌شوند:

$$E[g(X,Y)] = \sum_{Y \text{ برد}} \sum_{X \text{ برد}} g(X,Y) f_{X,Y}(x,y) \quad \text{برای متغیرهای گسسته } X \text{ و } Y$$

$$E[g(X,Y)] = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} g(X,Y) f_{X,Y}(x,y) dx dy \quad \text{برای متغیرهای پیوسته } X \text{ و } Y$$

تابع مولد گشتاورهای متغیر تصادفی (X,Y) :

$$m_{X,Y}(t_1, t_2) = E[e^{t_1 X + t_2 Y}]$$

توجه کنید که در این حالت نیز با گرفتن مشتق‌های پاره‌ای از تابع مولد گشتاور (X,Y) نسبت به t_1, t_2 و قرار دادن $(t_1, t_2) = (0,0)$ مقدار J آامین گشتاور توام (X,Y) بدست می‌آید. بصورت زیر:

$$\frac{\partial}{\partial t_1} m_{X,Y}(t_1, t_2) = \frac{\partial}{\partial t_1} E[e^{t_1 X + t_2 Y}] = E[X e^{t_1 X + t_2 Y}]$$

به همین ترتیب:

$$\frac{\partial}{\partial t_2} m_{X,Y}(t_1, t_2) = E[Y e^{t_1 X + t_2 Y}]$$

با ادامه دادن روند فوق در نهایت داریم:

$$\frac{\partial^{i+j}}{\partial t_1^i \partial t_2^j} m_{X,Y}(t_1, t_2) = E[X^i Y^j e^{t_1 X + t_2 Y}]$$

با قرار دادن $(t_1, t_2) = (0, 0)$ گشتاورهای (X, Y) بدست می‌آیند که عبارتند از:

$$m_{ij} = E[X^i Y^j] \quad i, j = 0, 1, 2, \dots$$

(i و j هر دو با هم صفر نمی‌باشند)

توجه کنید که m_{i0} و m_{0j} همان گشتاورهای متغیرهای تصادفی X و Y به تنهایی می‌باشند یعنی:

$$m_{i0} = E[X^i]$$

$$m_{0j} = E[Y^j]$$

همینطور:

$$m_{11} = E[X, Y]$$

$$m_{10} = E[X] \quad m_{02} = E[Y^2]$$

$$m_{01} = E[Y] \quad m_{21} = E[X^2 Y], \dots$$

۹-۵ مثال ۴: برای مثال ۲ مقدار مورد انتظار برای X ، Y و مجموع طول و عرض مشاهده شده را بدست بیاورید:

حل: طبق تعریف $E[X]$ بصورت زیر محاسبه می‌شود:

$$E[X] = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} x f_{X,Y}(x, y) dx dy = \int_0^1 \int_{-\sqrt{1-x^2}}^{\sqrt{1-x^2}} x \frac{1}{\pi} dy dx$$

$$= \int_0^1 \int_0^{2\pi} \frac{1}{\pi} r^2 \cos \theta d\theta dr = \frac{r^3}{3} \Big|_0^1 \sin \theta \int_0^{2\pi} = \frac{1}{\pi} \frac{1}{3} (0 - 0) = 0$$

البته از آنجا که بازه انتگرال گیری متقارن بوده و تابع نیز فرد می‌باشد می‌توانستیم بدون محاسبه انتگرال نیز صفر بودن جواب را بدست بیاوریم.

مقدار $E[X]$ نیز همانند $E[X]$ بصورت زیر بدست می‌آید:

$$E[X] = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} y f_{X,Y}(x, y) dx dy = \int_0^1 \int_{-\sqrt{1-y^2}}^{\sqrt{1-y^2}} y \frac{1}{\pi} dy dx = 0$$

باز هم به علت تقارن و فرد بودن تابع انتگرال فوق صفر می‌باشد.

۱۰-۵ حال مقدار $E[X+Y]$ را محاسبه می‌کنیم:

$$E[X+Y] = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} (X+Y) f_{X,Y}(x, y) dx dy = \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} [x f_{X,Y}(x, y) + y f_{X,Y}(x, y)] dx dy$$

$$= \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} x f_{X,Y}(x, y) dx + \int_{-\infty}^{+\infty} \int_{-\infty}^{+\infty} y f_{X,Y}(x, y) dy = E[X] + E[Y] = 0 + 0 = 0$$

توجه کنید که از مثال فوق می‌توان به این نتیجه رسید که $E[X+Y] = E[X] + E[Y]$ و در حالت کلی اگر $G(X)$ و $H(Y)$ به ترتیب توابعی از X و Y باشند داریم:

$$E[G(X) + H(Y)] = E[G(X)] + E[H(Y)]$$

همچنین توجه کنید که برای محاسبه $E[G(X)] + H(Y)$ نیازی به دانستن تابع چگالی توام متغیرهای تصادفی X و Y نمی‌باشد بلکه تنها با داشتن توابع چگالی کناری $f_X(x)$ و $f_Y(y)$ می‌توان آنرا محاسبه نمود.

۱۱-۵ کوواریانس یا همپراشی

برای مقایسه میزان وابستگی میان دو متغیر تصادفی X و Y از شاخصی به نام کوواریانس یا همپراشی X و Y استفاده می‌کنیم. که به صورت زیر تعریف می‌شود:

$$\delta_{X,Y} = \text{COV}(X, Y) = E[(X - \mu_x)(Y - \mu_y)]$$

توجه کنید که مقدار $\text{COV}(X, Y)$ در صورتی مثبت خواهد بود که اگر X بزرگتر از میانگینش باشد آنگاه Y نیز چنین باشد به عبارتی اگر X و Y هر دو هم جهت با یکدیگر افزایش یا کاهش داشته باشند مقدار کوواریانس مثبت خواهد بود

$$X \uparrow, Y \uparrow \Rightarrow \text{Sign}(\text{cov}(X, Y)) = +1$$

$$X \downarrow, Y \downarrow \Rightarrow \text{Sign}(\text{cov}(X, Y)) = +1$$

به همین ترتیب اگر $(X - \mu_x)$ و $(Y - \mu_y)$ با احتمال زیاد دارای علامت مخالف باشند یا به عبارتی X و Y در خلاف جهت هم باشند یعنی با

افزایش یکی دیگری کاهش پیدا کند یا بالعکس در این صورت کوواریانس X و Y منفی خواهد بود:

$$X \uparrow, Y \downarrow \Rightarrow \text{sign}(\text{cov}(X, Y)) = -1$$

$$X \downarrow, Y \uparrow \Rightarrow \text{sign}(\text{cov}(X, Y)) = -1$$

مقدار کوواریانس را از رابطه زیر که ساده‌تر می‌باشد نیز می‌توان بدست آورد:

$$\text{cov}(X, Y) = E[XY] - E[X]E[Y]$$

$$\text{cov}(X, Y) = E[(X - \mu_x)(Y - \mu_y)] = E[X Y - \mu_x Y - \mu_y X + \mu_x \mu_y]$$

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

اثبات:

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

۱۲-۵ با استفاده از رابطه فوق خواص متعددی را می‌توان برای کوواریانس بدست آورد که عبارتند از:

$$1- \text{cov}(X, Y) = \text{var}(X) = \delta_X^2$$

$$\text{cov}(X, Y) = E[X \cdot X] - E[X]E[X] = E[X^2] - E^2[X] = \delta_X^2$$

اثبات:

$$2- \text{cov}(X, Y) = \text{cov}(Y, X)$$

$$\text{cov}(X, Y) = E[XY] - E[X]E[Y] = E[YX] - E[Y]E[X] = \text{cov}(Y, X)$$

اثبات:

$$3- \text{cov}(X, Y) = 0 \quad (C \text{ مقدار ثابت})$$

$$\text{cov}(X, Y) = E[X \cdot C] - E[X]E[C] = CE[X] - CE[X] = 0$$

اثبات:

$$4- \text{cov}(ax \pm b, cy \pm d) = ac \text{ cov}(X, Y) \quad (a, b, c, d \text{ مقادیر اثبات})$$

$$\text{cov}(ax + b, cy + d) = E[(ax + b)(cy + d)] - E[ax + b]E[cy + d]$$

$$= E[acxy + adx + bcy + bd] - (aE[X] + b)(cE[Y] + d)$$

اثبات:

$$= acE[XY] + adE[X] + bcE[Y] + bd - acE[X]E[Y] - adE[X] - bcE[Y] - bd$$

$$a c E[XY] - a c E[X] E[Y] = a c (E[XY] - E[X] E[Y]) = a c \text{cov}(X, Y)$$

بنابراین تغییر مبدأ تأثیری روی مقدار کوواریانس ندارد اما تغییر مقیاس بر روی مقدار کوواریانس موثر می‌باشد.

$$\Delta - \text{cov}(X, Y \pm Z) = \text{cov}(X, Y) \pm \text{cov}(X, Z) \quad (X, Y, Z \text{ متغیرهای تصادفی})$$

$$\text{cov}(X, Y \pm Z) = E[X(Y \pm Z)] - E[X] E[Y \pm Z]$$

اثبات:

$$= E[XY \pm XZ] - E[X] [E[Y] \pm E[Z]]$$

$$= E[XY] \pm E[XZ] - E[X] [E[Y] \pm E[X] E[Z]]$$

$$= E[XY] - E[X] E[Y] \pm (E[XZ] \pm E[X] E[Z])$$

$$= \text{cov}(X, Y) \pm \text{cov}(X, Z)$$

۱۳-۵-۱۲.۴ ضریب همبستگی خطی بین X و Y

مقدار کوواریانس نیز همانند واریانس به مقیاس متغیرهای تصادفی X و Y وابسته است این مطلب را در خواص کوواریانس نیز نشان دادیم. به عبارتی ممکن است برای یک قانون احتمال، مقدار کوواریانس بسیار بیشتر از دیگری باشد، اما نمی‌توان گفت تمایل تغییر کردن X و Y با یکدیگر در حالت اول بیشتر از حالت دوم می‌باشد مشابه ایم مطلب را در فصل اول نیز برای واریانس دو جامعه آماری داشتیم که برای حل آن از ضریب تغییرات استفاده نمودیم. در اینجا برای بدست آوردن معیاری دقیق برای مقایسه میزان تمایل تغییرات همزمان دو متغیر C و Y از ضریب همبستگی خطی استفاده می‌کنیم که بصورت زیر تعریف می‌شود و با $f_{X,Y}$ نمایش داده می‌شود:

$$f_{X,Y} = \text{cov} \left(\frac{X - \mu_X}{\delta_X}, \frac{Y - \mu_Y}{\delta_Y} \right) = \frac{\delta_{XY}}{\delta_X \delta_Y}$$

همچنین با ساده نمودن رابطه فوق داریم:

$$\text{cov} \left(\frac{X - \mu_X}{\delta_X}, \frac{Y - \mu_Y}{\delta_Y} \right) = E \left[\left(\frac{X - \mu_X}{\delta_X} \right) \left(\frac{Y - \mu_Y}{\delta_Y} \right) \right] - E \left[\frac{X - \mu_X}{\delta_X} \right] E \left[\frac{Y - \mu_Y}{\delta_Y} \right]$$

$$E \left[\frac{X - \mu_X}{\delta_X} \right] = E \left[\frac{Y - \mu_Y}{\delta_Y} \right] = 0 \quad \text{برمال شده می‌باشد بنابراین امید ریاضی آن برابر صفر می‌باشد یعنی:}$$

$$f_{X,Y} = E \left[\left(\frac{X - \mu_X}{\delta_X} \right) \left(\frac{Y - \mu_Y}{\delta_Y} \right) \right] \quad \text{بنابراین داریم:}$$

۱۴-۵-۱۲-۵ حال نشان می‌دهیم که مقدار $f_{X,Y}$ همواره بین ۱ و -۱ قرار دارد یعنی $|f_{X,Y}| \leq 1$

اثبات: دو متغیر تصادفی W و Z و مقدار متغیر a را در نظر بگیرید در این صورت متغیر تصادفی $T = (aW - Z)^2$ را تعریف می‌کنیم. به ازای هر مقادیر متغیر تصادفی T همواره مثبت است بنابراین امید ریاضی آنها نیز مثبت می‌باشد یعنی داریم: $E[(aW - Z)^2] \geq 0$

$$\Rightarrow E[a^2 W^2 - 2aWZ + Z^2] \geq 0$$

$$\Rightarrow a^2 E[W^2] - 2a E[WZ] + E[Z^2] \geq 0 \quad ; \forall a$$

$$a = \frac{E[WZ]}{E[W^2]} \quad \text{حال چون به ازای هر } a \text{ نامساوی فوق برقرار می‌باشد می‌توان قرار داد}$$

بنابراین:

$$\Rightarrow \frac{E[wz]}{E[w^2]} E[w^2] - \frac{E[wz]^2}{E[w^2]} + E[z^2] \geq 0$$

$$\Rightarrow -\frac{E[wz]^2}{E[w^2]} + E[z^2] \geq 0 \Rightarrow \frac{E[wz]^2}{E[w^2] E[z^2]} \leq 1$$

حال به جای متغیرهای تصادفی w و z قرار می‌دهیم: $w = X - \mu_X$, $z = Y - \mu_Y$

$$\frac{E\left[\frac{(X-\mu_X)(Y-\mu_Y)}{X} \right]}{E\left[\frac{(X-\mu_X)^2}{X}\right] E\left[\frac{(Y-\mu_Y)^2}{Y}\right]} \leq 1 \Rightarrow \frac{\delta_{XY}^2}{\delta_X^2 \delta_Y^2} \leq 1$$

داریم:

$$\Rightarrow \rho_{XY}^2 \leq 1 \Rightarrow |\rho_{XY}| \leq 1$$

۱۵-۵ توجه کنید که در اثبات فوق اگر حالت تساوی در نامساوی $E[(aw-z)^2]$ رخ دهد داریم:

اما متغیر تصادفی T همواره مثبت است بنابراین میانگین آن تنها زمانی صفر می‌باشد که تمام مقادیری که T قبول می‌کند برابر صفر باشند به عبارت دقیقتر متغیر تصادفی T با احتمال ۱ برابر صفر می‌باشد یعنی:

$$\begin{aligned} p(T=0) &= 1 \Rightarrow p((aw-z)^2 = 0) = 1 \\ &\Rightarrow p(aw-z = 0) = 1 \\ &\Rightarrow p(aw=z) = 1 \\ &\Rightarrow p(Y-\mu_Y = a(X-\mu_X)) = 1 \\ &\Rightarrow p(Y = ax + \mu_Y - a\mu_X) = 1 \end{aligned}$$

۱۶-۵ بنابراین با احتمال صددرصد متغیر تصادفی Y برابر با $ax + \mu_Y - a\mu_X$ می‌باشد و این نشان می‌دهد که از مقادیر قابل استفاده به (X, Y)

تنها مقادیری که روی یک خط راست به معادله $Y = ax + \mu_Y - a\mu_X$ قرار دارند می‌توانند احتمال مثبت داشته باشد. به همین دلیل است که

می‌گوییم ضریب همبستگی خطی میزان وابستگی خطی بین دو متغیر X و Y را نشان می‌دهد.

هر چه مقدار $P_{X,Y}$ به عدد ۱ یا -۱ نزدیک‌تر باشد، متغیرهای X و Y تمایل بیشتری به متغیر مستقیم یا معکوس با یکدیگر خواهند داشت و اگر $P_{X,Y} = 0$ باشد، بدین معنی است که متغیرهای X و Y به صورت خطی به یکدیگر وابسته نمی‌باشند. مثال زیر نشان می‌دهد که ممکن است برای دو متغیر X و Y داشته باشیم $P_{X,Y} = 0$ اما در عین حال دو متغیر کاملاً به یکدیگر وابسته باشند.

۱۷-۵ مثال ۵: متغیر تصادفی X بصورت زیر تعریف شده است:

$$f_X(x) = \begin{cases} \frac{1}{4} & -2 \leq x \leq 2 \\ 0 & \text{سایر مقادیر} \end{cases}$$

متغیر تصادفی Y را برابر X^2 تعریف می‌کنیم در این صورت:

الف) $f_Y(y)$ را محاسبه کنید.

ب) ρ_{XY} را محاسبه کنید.

حل: با توجه به اینکه $Y = X^2$ داریم:

$$F_Y(y) = p(Y \leq y) = p(X^2 \leq y) = p(-\sqrt{y} \leq X \leq \sqrt{y}) = F_X(\sqrt{y}) - F_X(-\sqrt{y})$$

حال می‌بایستی تابع توزیع متغیر تصادفی X را بدست بیاوریم:

$$F_X(x) = \int_{-\infty}^x f_X(x) dx = \int_{-2}^x \frac{1}{4} dx = \frac{1}{4}(x+2)$$

○ $x \leq -2$

$$F_X(x) = \begin{cases} \frac{1}{4}(x+2) & -2 \leq x \leq 2 \\ 1 & x > 2 \end{cases}$$

بنابراین تابع توزیع Y برابر است با:

$$F_Y(y) = \frac{1}{4}(\sqrt{y}+2) - \frac{1}{4}(-\sqrt{y}+2) = \frac{1}{2}\sqrt{y}$$

$$F_Y(y) = \begin{cases} 0 & y < 0 \\ \frac{\sqrt{y}}{2} & 0 \leq y \leq 4 \\ 1 & 4 < y \end{cases}$$

با مشتق‌گیری از $F_Y(y)$ تابع چگالی Y را بدست می‌آوریم.

$$f_Y(y) = \begin{cases} \frac{1}{4\sqrt{y}} & 0 \leq y \leq 4 \\ 0 & \text{سایر} \end{cases}$$

۵-۱۸ ب: برای محاسبه ρ_{XY} به مقادیر $E[XY]$ ، $E[Y]$ ، $E[X]$ نیاز داریم زیرا:

$$\rho_{XY} = \frac{\delta_{XY}}{\delta_X \delta_Y} = \frac{E[XY] - E[X]E[Y]}{\delta_X \delta_Y}$$

$$E[X] = \int_{-\infty}^{+\infty} x f_X(x) dx = \int_{-2}^2 \frac{1}{4} x dx = 0$$

به دلیل فرد بودن تابع x در بازه $[-2, 2]$ مقدار انتگرال صفر می‌باشد.

$$E[Y] = \int_0^4 y \frac{1}{4} \frac{1}{\sqrt{y}} dy = \frac{1}{4} \left(\frac{2}{3} y^{\frac{3}{2}} \right) \Big|_0^4 = \frac{8}{6}$$

$$E[XY] = E[XX^2] = E[X^3] = \int_{-2}^2 x^3 \frac{1}{4} dx = 0$$

به این ترتیب مقدار ρ_{XY} برابر است با:

$$\delta_{XY} = 0 - \frac{8}{6} \times 0 = 0 \Rightarrow \rho_{XY} = 0$$

کوواریانس اول

همانطور که ملاحظه می‌کنید مقدار ρ_{XY} صفر می‌باشد و این به معنی عدم وابستگی دو متغیر X و Y نمی‌باشد بلکه به معنی عدم وابستگی خطی آن‌هاست به وضوح X و Y بصورت توانی ($Y = X^2$) به یکدیگر وابسته‌اند.

مثال ۶: نشان دهید ضریب همبستگی خطی با اعمال تغییر مبدأ و مقیاس روی متغیرهای تصادفی X و Y بدون تغییر باقی می‌ماند؟

حل: برای این منظور می‌بایستی ثابت کنیم:

$$\rho_{ax+b, cy+d} = \rho_{XY}$$

اثبات:

$$\rho_{ax+b, cy+d} = \frac{\text{cov}(ax+b, cy+d)}{\sqrt{\text{var}(ax+b) \text{var}(cy+d)}} = \frac{ac \text{cov}(X, Y)}{\sqrt{a^2 \text{var}(X) c^2 \text{var}(Y)}}$$

$$= \frac{\text{cov}(X, Y)}{\sqrt{\text{var}(X) \text{var}(Y)}} = \rho_{X, Y}$$

۲۰-۵ مثال ۷: در مثال ۱ مقدار $\text{var}(3x - 2y + 1)$ را محاسبه کنید:

در مثال ۱ جدول احتمالات دو متغیره بصورت زیر بدست آمد:

$Y \backslash X$	۲	۳	۴	
-۱	۰	$\frac{1}{4}$	۰	$\frac{1}{4}$
۰	$\frac{1}{4}$	۰	$\frac{1}{4}$	$\frac{2}{4}$
۱	۰	$\frac{1}{4}$	۰	$\frac{1}{4}$
	$\frac{1}{4}$	$\frac{1}{2}$	$\frac{1}{4}$	

$$\text{var}(X) = \text{cov}(X, X) \quad \text{می‌دانیم:}$$

بنابراین می‌توان نوشت:

$$\text{var}(3X - 2Y + 1) = \text{var}(3X - 2Y) = \text{cov}(3X - 2Y, 3X - 2Y)$$

$$= \text{cov}(3X, 3X) + \text{cov}(2Y, 2Y) - 12 \text{cov}(X, Y)$$

$$= 9 \text{var}(X) + 4 \text{var}(Y) - 12 \text{cov}(X, Y)$$

بنابراین می‌بایستی مقادیر δ_X^2 , δ_Y^2 , δ_{XY} را از روی جدول محاسبه کنیم:

$$E[X] = \frac{1}{4} \times 2 + \frac{1}{2} \times 3 + \frac{1}{4} \times 4 = \frac{1}{2} + \frac{3}{2} + 1 = 3$$

$$E[Y] = \frac{1}{4}(-1) + \frac{2}{4} \times 0 + \frac{1}{4} \times 1 = 0$$

$$E[X^2] = \frac{1}{4} \times 4 + \frac{1}{2} \times 9 + \frac{1}{4} \times 16 = 1 + \frac{9}{2} + 4 = \frac{19}{2}$$

$$E[Y^2] = \frac{1}{4} \times 1 + \frac{2}{4} \times 0 + \frac{1}{4} \times 1 = \frac{1}{2}$$

$$\text{var}(X) = E[X^2] - E^2[X] = \frac{19}{2} - 9 = \frac{1}{2}$$

$$\text{var}(Y) = E[Y^2] - E^2[Y] = \frac{1}{2} - 0 = \frac{1}{2}$$

$$E[XY] = (3 \times -1) \times \frac{1}{4} + (3 \times 1) \times \frac{1}{4} \times (4 \times 0) \times \frac{1}{4} \times (3 \times 0) \times \frac{1}{4} = 0$$

$$\Rightarrow \text{cov}(X, Y) = E[XY] - E[X]E[Y] = 0$$

$$\Rightarrow \text{var}(3X - 2Y + 1) = 9 \times \frac{1}{2} + 4 \times \frac{1}{2} - 0 = \frac{13}{2} = 6.5$$

فصل ششم

۶-۱-۱

۱.۶ برخی توابع توزیع گسسته

در فصل قبل با چگونگی بدست آوردن توابع توزیع و چگالی آشنا شدید در این فصل قصد داریم چند تابع چگالی و توزیع خاص را مورد بررسی قرار دهیم. بسیاری از آزمایش‌های تصادفی با وجود اینکه ظاهراً متفاوت می‌باشند اما از ماهیت یکسانی برخوردار می‌باشند و از یک الگو پیروی می‌کنند به همین دلیل است که بررسی توابع چگالی و توزیع آنها مهم می‌باشد.

۱.۱.۶ متغیر تصادفی برنولی و دو جمله‌ای

اگر برای یک آزمایش تصادفی تنها دو نتیجه موفقیت و شکست امکان‌پذیر باشد به آن آزمایش، آزمایش برنولی می‌گوییم و احتمال موفقیت را با p و شکست را با q نمایش می‌دهیم. مثلاً پرتاب یک سکه یک آزمایش برنولی است که دو حالت ممکن را در پی دارد. می‌توانیم آمدن شیر را موفقیت و آمدن خط را شکست در نظر بگیریم. بدیهیست که در این حالت $p = \frac{1}{2}$, $q = 1 - p = \frac{1}{2}$.

۶-۱-۲

یک متغیر تصادفی برنولی عبارتست از تعداد پیروزی (۰ یا ۱ بار) در یک مرتبه انجام آزمایش برنولی. به این ترتیب تابع چگالی برای متغیر تصادفی برنولی X بصورت زیر خواهد بود:

$$f_X(x) = \begin{cases} p & x=1 \\ q & x=0 \\ 0 & \text{سایر مقادیر} \end{cases}$$

متغیر تصادفی برنولی را با نماد $B(1, p)$ نمایش می‌دهیم که در آن عدد ۱ نمایش دهنده یکبار انجام آزمایش است. تابع چگالی متغیر تصادفی برنولی را بصورت زیر هم می‌توان نوشت:

$$f(x) = p^x q^{1-x} \quad x=0, 1 \\ = 0 \quad \text{سایر مقادیر}$$

حال امید، واریانس و تابع مولد گشتاور متغیر تصادفی برنولی را بدست می‌آوریم:

$$E[X] = 1 \times p + 0 \times q = p$$

$$E[X^2] = 1^2 \times p + 0^2 \times q = p$$

$$\text{var}(x) = E[x^2] - E^2[X] = p - p^2 = p(1-p) = pq$$

پس:

به همین ترتیب:

$$m_X(t) = E[e^{tx}] = pe^{t \times 1} + qe^{t \times 0} = pe^t + q$$

۶-۲

مثال ۱: شخصی ۵ بلیط می‌فروشد که یکی از آنها قلابی است از وی یک بلیط خریداری می‌نیم، قرار می‌دهیم:

$$X = \begin{cases} \text{اگر بلیط اصلی را خریداری کرده باشیم} & 1 \\ \text{اگر بلیط قلابی را خریداری کرده باشیم} & 0 \end{cases}$$

در این صورت تابع چگالی متغیر تصادفی X را بدست آورید. و $E[X]$, $\text{var}[X]$ را بدست آورید.

حل: با توجه به اینکه متغیر تصادفی X دو حالت موفقیت و شکست را نشان می‌دهد بنابراین از متغیر تصادفی برنولی پیروی می‌کند. حالت مقدار احتمال p و q را محاسبه می‌کنیم:

$$p = \frac{4}{5} = \text{احتمال اینکه بلیط اصلی خریداری شود}$$

$$q = 1 - p = \frac{1}{5}$$

بنابراین تابع چگالی برابر است با:

$$f(x) = \begin{cases} \frac{4}{5} & x=1 \\ \frac{1}{5} & x=0 \\ 0 & \text{سایر مقادیر} \end{cases}$$

$$E[X] = p = \frac{4}{5}$$

$$\text{var}[X] = pq = \frac{4}{5} \frac{1}{5} = \frac{4}{25}$$

عدد $E[X] = \frac{4}{5}$ نشان می‌دهد که اگر از شخصی مورد نظر مثلاً ۱۰۰ عدد بلیط خریداری کنیم باید انتظار داشته باشیم $100 \times \frac{4}{5} = 80$ عدد از بلیط‌ها اصلی باشند.

۳-۶ متغیر تصادفی دو جمله‌ای

اگر یک آزمایش برنولی را n بار بطور مستقل انجام دهیم و متغیر تصادفی X را برابر با تعداد پیروزی‌ها در این n بار انجام آزمایش در نظر بگیریم، در این صورت متغیر تصادفی X را دو جمله‌ای می‌نامیم.

برای بدست آوردن تابع چگالی متغیر تصادفی دو جمله‌ای می‌بایستی مقادیری که X قبول می‌کند و مقدار احتمال آنرا بدست بیاوریم. از آنجا که متغیر تصادفی دو جمله‌ای تعداد پیروزی‌ها در n بار از انجام مستقل آزمایش برنولی می‌باشد بنابراین تعداد پیروزی‌ها می‌تواند هر یک از مقادیر 0 تا n باشد و احتمال اینکه x بار پیروز و در $n-x$ بار آزمایش باقیمانده شکست بخوریم با توجه به استقلال آزمایشها برابر است با:

$$(p \cdot \underbrace{p \dots p}_x) (q \cdot \underbrace{q \dots q}_{n-x}) = p^x q^{n-x}$$

اما این مساله که در کدام یک از n آزمایش پیروز شویم نیز مهم است. با توجه به قواعد شمارشی به $\binom{n}{x}$ صریق می‌توان در n آزمایش پیروز شد بنابراین تابع چگالی احتمال برای متغیر تصادفی دو جمله‌ای X برابر است با:

$$f_X(x) = \binom{n}{x} p^x q^{n-x} \quad x = 0, 1, 2, \dots, n$$

$$= 0 \quad \text{سایر مقادیر}$$

اگر $X_1, X_2, \dots, X_n$ مجموعه n متغیر تصادفی برنولی باشند در این صورت با توجه به تعریف متغیر تصادفی دو جمله‌ای X عبارتست از:

$$X = X_1 + X_2 + \dots + X_n$$

متغیر تصادفی دو جمله‌ای را بصورت $X \sim B(n, p)$ نمایش می‌دهیم که شامل دو پارامتر می‌باشد:

p = احتمال پیروزی

n = تعداد آزمایشها

به این ترتیب متغیر تصادفی برنولی حالت خاصی از متغیر تصادفی دو جمله‌ای با پارامتر $n=1$ می‌باشد.

حال به محاسبه مقادیر امید ریاضی، واریانس و تابع مولد گشتاور برای متغیر تصادفی دو جمله‌ای می‌پردازیم.

برای محاسبه $E[X]$ توجه می‌کنیم که متغیر تصادفی دو جمله‌ای X برابر است با مجموع نتایج حاصله از n بار انجام مستقل آزمایش برنولی بنابراین:

$$E[X] = E\left[\sum_{i=1}^n X_i\right] = \sum_{i=1}^n E[X_i] = \sum_{i=1}^n p = np$$

$$\text{var}(X) = \text{VAR}\left(\sum_{i=1}^n X_i\right) = \sum_{i=1}^n \text{var}(X_i) = \sum_{i=1}^n pq = npq$$

به همین ترتیب:

$$\begin{aligned} m_X(t) &= E\left[e^{tx}\right] = E\left[e^{t\left(\sum_{i=1}^n X_i\right)}\right] = E\left[e^{tX_1 + tX_2 + \dots + tX_n}\right] \\ &= E\left[e^{tX_1} e^{tX_2} \dots e^{tX_n}\right] = E\left[e^{tX_1}\right] E\left[e^{tX_2}\right] \dots E\left[e^{tX_n}\right] \\ &= (pe^t + q)(pe^t + q) \dots (pe^t + q) = (pe^t + q)^n \end{aligned}$$

بار n

۵-۶ مثال ۲: اگر در مثال ۱ تعداد ۱۰ بلیط از فروشنده خریداری می‌کنیم و متغیر X را تعداد بلیط‌های اصلی در نظر بگیریم. مطلوبست:

الف) تابع چگالی احتمال متغیر تصادفی X .

ب) رسم نمودار تابع چگالی X .

ج) مقدار مد.

د) امید و واریانس متغیر X .

حل: الف) در این مساله می‌بایستی یک آزمایش برنولی را ۱۰ مرتبه بصورت مستقل تکرار کنیم بنابراین X یک متغیر تصادفی دو جمله‌ای می‌باشد و تابع چگالی آن عبارتست از:

$$X \sim B\left(10, \frac{4}{5}\right)$$

$$f_X(x) = \binom{10}{x} \left(\frac{4}{5}\right)^x \left(\frac{1}{5}\right)^{10-x} \quad x=0, 1, 2, \dots, 10$$

$= 0$ سایر مقادیر

حال مقدار احتمال را به ازای $x=0, \dots, 10$ بدست می‌آوریم:

$$f_X(x=0) = \binom{10}{0} \left(\frac{4}{5}\right)^0 \left(\frac{1}{5}\right)^{10} = \frac{1}{5^{10}} = 0$$

$$f_X(x=1) = \binom{10}{1} \left(\frac{4}{5}\right)^1 \left(\frac{1}{5}\right)^9 = 10 \cdot \frac{4}{5^{10}} = 0$$

$$f_X(x=2) = \binom{10}{2} \left(\frac{4}{5}\right)^2 \left(\frac{1}{5}\right)^8 = 45 \cdot \frac{4^2}{5^{10}} = 0$$

$$f_X(x=3) = \binom{10}{3} \left(\frac{4}{5}\right)^3 \left(\frac{1}{5}\right)^7 = 120 \cdot \frac{4^3}{5^{10}} = 0/\dots 7$$

$$f_X(x=3) = \binom{10}{3} \left(\frac{4}{5}\right)^3 \left(\frac{1}{5}\right)^7 = 120 \cdot \frac{4^3}{5^{10}} = 0/\dots 7$$

$$f_X(x=4) = \binom{10}{4} \left(\frac{4}{5}\right)^4 \left(\frac{1}{5}\right)^6 = 210 \cdot \frac{4^4}{5^{10}} = 0.0055$$

$$f_X(x=5) = \binom{10}{5} \left(\frac{4}{5}\right)^5 \left(\frac{1}{5}\right)^5 = 252 \cdot \frac{4^5}{5^{10}} = 0.026$$

$$f_X(x=6) = \binom{10}{6} \left(\frac{4}{5}\right)^6 \left(\frac{1}{5}\right)^4 = 210 \cdot \frac{4^6}{5^{10}} = 0.088$$

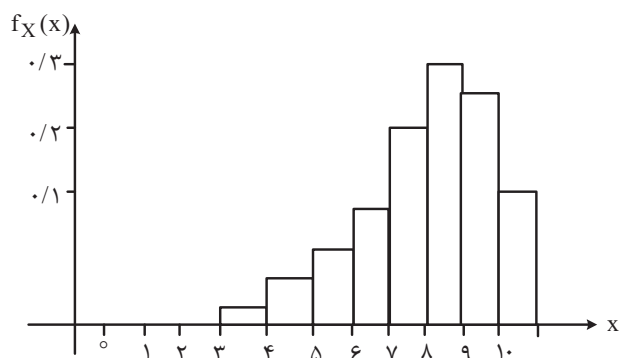
$$f_X(x=7) = \binom{10}{7} \left(\frac{4}{5}\right)^7 \left(\frac{1}{5}\right)^3 = 120 \cdot \frac{4^7}{5^{10}} = 0.2$$

$$f_X(x=8) = \binom{10}{8} \left(\frac{4}{5}\right)^8 \left(\frac{1}{5}\right)^2 = 45 \cdot \frac{4^8}{5^{10}} = 0.3$$

$$f_X(x=9) = \binom{10}{9} \left(\frac{4}{5}\right)^9 \left(\frac{1}{5}\right)^1 = 10 \cdot \frac{4^9}{5^{10}} = 0.268$$

$$f_X(x=10) = \binom{10}{10} \left(\frac{4}{5}\right)^{10} \left(\frac{1}{5}\right)^0 = 1 \cdot \frac{4^{10}}{5^{10}} = 0.1$$

(ب) حال نمودار مستطیلی تابع چگالی متغیر تصادفی X بصورت زیر بدست می‌آید:



توجه کنید که $\sum_{x=0}^{10} \binom{10}{x} \left(\frac{4}{5}\right)^x \left(\frac{1}{5}\right)^{10-x} = 1$ می‌باشد.

(ج) از فصل دوم به یاد دارید که مد مقداری از X است که به ازای آن تابع چگالی ماکزیمم می‌شود. از روی نمودار به وضوح پیداست که مد متغیر تصادفی X برابر با: $x=8$ می‌باشد. یعنی به ازای هر بار خرید ۱۰ عدد بلیط از فروشنده به احتمال زیاد (۰/۳) ۸ میانه توزیع نیز می‌باشد زیرا

$$F_X(x=8) \geq \frac{1}{2} \text{ می‌شود یعنی: } \frac{1}{2} \text{ بزرگتر یا مساوی با } \frac{1}{2} \text{ می‌شود یعنی: } F_X(x=8) \geq \frac{1}{2}$$

(د)

$$E[X] = np = 10 \times \frac{4}{5} = 8$$

$$\text{var}(X) = npq = 10 \times \frac{4}{5} \times \frac{1}{5} = \frac{8}{5}$$

همانطور که ملاحظه می‌کنید مقدار امید ریاضی X برابر ۸ می‌باشد به همین دلیل است که در مثال ۱ نشان دادیم که اگر ۱۰۰ بلیط از فروشنده

دریافت کنیم به طور متوسط $100 \times \frac{4}{5} = 80$ عدد از آنها اصلی می‌باشند.

۱-۷-۶ ۱.۲.۶ مد توزیع دو جمله‌ای

در مثال قبل برای بدست آوردن مد از تابع چگالی استفاده نمودیم اما می‌توان مقدار مد را بدون استفاده از تابع چگالی و تنها با داشتن مقادیر پارامترهای توزیع دو جمله‌ای بدست آورد. می‌دانیم متغیر تصادفی دو جمله‌ای گسسته می‌باشد از طرفی چون تابع چگالی به ازای مقدار مد بیشترین مقدار را قبول می‌کند بنابراین داریم:

$$\begin{cases} f(x_M) \geq f(x_M - 1) & (1) \\ f(x_M) \geq f(x_M + 1) & (2) \end{cases}$$

$$(1) \Rightarrow \binom{n}{x_M} p^{x_M} q^{n-x_M} \geq \binom{n}{x_M+1} p^{x_M+1} q^{n-(x_M+1)}$$

$$\Rightarrow \frac{n!}{x_M!(n-x_M)!} \geq \frac{n!}{(x_M+1)!(n-(x_M+1))!} p q^{-1}$$

$$x_M + 1 \geq (n - x_M) p q^{-1} \Rightarrow x_M + x_M p q^{-1} \geq n p q^{-1} - 1$$

$$\Rightarrow x_M \geq n p^{-1} + p \Rightarrow x_M \geq p(n+1) - 1$$

۲-۷-۶ به همین ترتیب از رابطه (۲) می‌توان نتیجه گرفت:

$$x_M \leq (n+1)p$$

در نتیجه داریم:

$$p(n+1) - 1 \leq x_M \leq p(n+1)$$

حال اگر $p(n+1)$ عددی صحیح باشد $p(n+1) - 1$ نیز عددی صحیح است و هر دو مد متغیر تصادفی X می‌باشند در غیر این صورت بین دو عدد $p(n+1)$ و $p(n+1) - 1$ تنها یک عدد صحیح وجود دارد که آن عدد مد توزیع دو جمله‌ای می‌باشد.

مثال ۳: مد تغییر تصادفی دو جمله‌ای در مثال قبل را بدون استفاده از تابع چگالی بدست آورید.

حل: X دارای توزیع دو جمله‌ای با پارامترهای $n=10$ ، $p=\frac{4}{5}$ می‌باشد.

بنابراین:

$$\frac{4}{5}(10+1) - 1 \leq x_M \leq \frac{4}{5}(10+1)$$

$$7 + \frac{1}{5} \leq x_M \leq 8 + \frac{1}{5} \Rightarrow x_M = 8$$

۱-۸-۶ ۳.۶ متغیر تصادفی هندسی (نوع اول)

اگر متغیر تصادفی X برابر باشد با تعداد شکست‌ها قبل از رسیدن به اولین پیروزی در انجام آزمایشهای مستقل برنولی، در این صورت به آن متغیر تصادفی هندسی از نوع اول می‌گوییم. به عنوان مثال تعداد دفعات پرتاب یک سکه برای بدست آوردن اولین شیر یا تعداد دفعات شلیک به هدف تا قبل از اولین برخورد گلوله با هدف، نمونه‌هایی از متغیرهای تصادفی هندسی نوع اول می‌باشند.

برای بدست آوردن تابع چگالی متغیر تصادفی هندسی نوع اول توجه می‌کنیم که X مقادیر 0 ، 1 ، $\dots$ را قبول می‌کند و از آنجا که در آخرین آزمایش پیروز می‌شویم بنابراین در X آزمایش قبل شکست خورده‌ایم که به دلیل استقلال آزمایشها، احتمال آن برابر است q^X و برای اینکه در آزمایش بعدی پیروز شویم می‌بایستی احتمال p (پیروزی) را در عدد q^X ضرب کنیم بنابراین:

$$f_X(x) = p q^x \quad x=0, 1, \dots$$

○ سایر مقادیر

حال امید ریاضی، واریانس و تابع مولد گشتاور متغیر تصادفی هندسی نوع اول را محاسبه می‌کنیم:

$$M_X(t) = E[e^{tX}] = \sum_{x=0}^{\infty} e^{tx} p q^x = \sum_{x=0}^{\infty} e^{tx} p (q e^t)^x = \frac{p}{1 - q e^t}$$

برای بدست آوردن امید ریاضی و واریانس از تابع مولد گشتاور استفاده می‌کنیم:

$$E[X] = M'_X(t) |_{t=0} = \frac{p q e^t}{(1 - q e^t)^2} |_{t=0} = \frac{p q}{(1 - q)^2} = \frac{q}{p}$$

$$E[X^2] = M''_X(t) |_{t=0} = \frac{p q e^t (1 - q e^t)^{-2} + 2 q e^t (1 - q e^t)^{-3} p q e^t}{(1 - q e^t)^4} |_{t=0} = \frac{q p + q^2}{p^2}$$

پس:

$$\text{var}(X) = E[X^2] - E^2[X] = \frac{p q + q^2}{p^2} - \frac{q^2}{p^2} = \frac{q(p + q)}{p^2} = \frac{q}{p^2}$$

۹-۶ مثال ۴: شخصی به هدفی شلیک می‌کند. گلوله با احتمال $\frac{3}{5}$ به هدف برخورد می‌کند اگر متغیر تصادفی X را تعداد شلیکها قبل از مورد

اصابت قرار دادن هدف در نظر بگیریم. مطلوبست:

الف) تابع چگالی متغیر تصادفی X .

ب) احتمال اینکه در شلیک ۷ام هدف را بزند.

ج) به طور متوسط چند بار شلیک قبل از اینکه هدف مورد اصابت قرار بگیرد لازم است.

د) واریانس و تابع مولد گشتاور را برای متغیر تصادفی X محاسبه کنید.

حل: الف) متغیر تصادفی X هندسی از نوع اول می‌باشد و با پارامتر $p = \frac{3}{5}$ بنابراین تابع چگالی آن بصورت زیر بدست می‌آید:

$$f(x) = \frac{3}{5} \left(\frac{2}{5}\right)^x \quad x = 0, 1, \dots$$

سایر مقادیر

ب) برای اینکه در شلیک ۷ام هدف را بزند می‌بایستی ۶ شلیک ناموفق داشته باشیم بنابراین:

ج) برای بدست آوردن متوسط تعداد شلیکها قبل از زدن بایستی امید ریاضی متغیر تصادفی X را بدست آوریم:

$$E[X] = \frac{q}{p} = \frac{\frac{2}{5}}{\frac{3}{5}} = \frac{2}{3}$$

بنابراین شخص می‌بایستی بیشتر اوقات پس از یک شکست هدف را بزند.

د)

$$\text{var}(X) = \frac{q}{p^2} = \frac{\frac{2}{5}}{\left(\frac{3}{5}\right)^2} = \frac{10}{9}$$

$$M_X(t) = \frac{p}{1 - q e^t} = \frac{\frac{3}{5}}{1 - \frac{2}{5} e^t} = \frac{3}{5 - 2 e^t}$$

متغیر تصادفی X را تعداد آزمایشها برای رسیدن به اولین پیروزی در آزمایشهای مستقل برنولی در نظر می‌گیریم. در این صورت X یک متغیر تصادفی هندسی نوع دوم است. برای بدست آوردن تابع چگالی متغیر تصادفی هندسی نوع دوم ابتدا توجه می‌کنیم که حداقل یک آزمایش برای رسیدن به اولین پیروزی مورد نیاز است، احتمال اینکه در آزمایش x ام پیروز شویم برابر است با احتمال اینکه در $x-1$ آزمایش قبلی شکست خورده و در آزمایش آخر پیروز شویم که اولی با احتمال q^{x-1} و دومی با احتمال p رخ می‌دهد. نتیجه تابع چگالی بصورت زیر بدست می‌آید:

$$f_X(x) = pq^{x-1} \quad x = 1, 2, \dots$$

= 0 سایر مقادیر

توجه کنید که بین متغیر تصادفی هندسی نوع اول و دوم رابطه مستقیمی برقرار است، از آنجا که متغیر تصادفی هندسی نوع دوم تعداد آزمایشها برای رسیدن به اولین پیروزی است بنابراین می‌بایستی در تمام آزمایشها شکست بخوریم غیر آزمایش آخر. اگر Y یک متغیر هندسی نوع دوم باشد و X یک متغیر هندسی نوع اول، رابطه $Y = X + 1$ بین ایندو برقرار می‌باشد. با کمک این رابطه می‌توان به راحتی امید، واریانس و تابع مولد گشتاور متغیر هندسی نوع دوم را از روی نوع اول بدست آورد:

$$E[Y] = E[X+1] = E[X] + 1 = \frac{q}{p} + 1 = \frac{1}{p}$$

$$\text{var}(Y) = \text{var}(X+1) = \text{var}(X) = \frac{q}{p^2}$$

$$M_Y(t) = E[e^{tY}] = E[e^{tY+t}] = e^t E[e^{tX}] = e^t \frac{p}{1-qe^t}$$

مثال ۵: متغیر تصادفی Y را تعداد شلیکها برای رسیدن به اولین اصابت گلوله به هدف در مثال قبل در نظر می‌گیریم:
مطلوبست:

الف) تابع چگالی متغیر تصادفی X .

ب) احتمال اینکه در کمتر یا مساوی با دئ آزمایش هدف مورد اصابت گلوله قرار گیرد.

ج) متوسط تعداد شلیکها برای زدن هدف.

د) واریانس و تابع مولد گشتاور متغیر تصادفی Y .

حل: الف) یک متغیر تصادفی هندسی نوع دوم است بنابراین تابع چگالی آن بصورت زیر می‌باشد:

$$f_Y(y) = \frac{3}{5} \left(\frac{2}{5}\right)^{y-1} \quad y = 1, 2, \dots$$

= 0 سایر مقادیر

ب) برای اینکه در کمتر یا مساوی با دو آزمایش هدف مورد اصابت قرار گیرد می‌بایستی احتمال $f(y=1) + f(y=2)$ را محاسبه کنیم:

$$f(y=1) + f(y=2) = \left(\frac{3}{5}\right) \left(\frac{2}{5}\right)^0 + \left(\frac{3}{5}\right) \left(\frac{2}{5}\right)^1 = \frac{3}{5} \left(1 + \frac{2}{5}\right) = \frac{21}{25} = 0.84$$

$$E[Y] = \frac{1}{p} = \frac{1}{\frac{3}{5}} = \frac{5}{3} = 1.66 \quad \text{ج)}$$

به همین ترتیب:

$$\text{var}(Y) = \frac{q}{p^2} = \frac{\frac{2}{5}}{\frac{9}{25}} = \frac{10}{9}$$

(د)

$$m_Y(t) = \frac{\frac{2}{5} e^t}{1 - \frac{2}{5} e^t} = \frac{2 e^t}{5 - 2 e^t}$$

۱۲-۶ ۴.۶ متغیر تصادفی دو جمله‌ای منفی (نوع اول)

اگر متغیر تصادفی X تعداد شکستها قبل از رسیدن به k امین پیروزی در تکرار مستقل آزمایشهای برنولی در نظر گرفته شود به آن متغیر تصادفی دو جمله‌ای منفی از نوع اول می‌گوییم.

توجه کنید که اگر $k=1$ باشد متغیر تصادفی دو جمله‌ای منفی نوع اول همان متغیر تصادفی هندسی نوع اول می‌باشد. برای بدست آوردن تابع چگالی متغیر تصادفی دو جمله‌ای منفی نوع اول توجه کنید که در آزمایش آخر پیروزی k ام رخ می‌دهد. بنابراین در آزمایشهای قبلی $k-1$ بار پیروزی و X بار شکست رخ داده است یعنی از بین $x+k-1$ آزمایش می‌بایستی $k-1$ بار پیروز شویم که احتمال آن p^{k-1} می‌باشد و احتمال اینکه X بار شکست بخوریم هم مهم است که به $\binom{x+k-1}{x}$ طریق ممکن است. نهایتاً در آخرین آزمایش هم با احتمال p پیروز می‌شویم که همان k امین پیروزی است. به همین ترتیب با ضرب موارد فوق تابع چگالی بصورت زیر بدست می‌آید:

$$f_X(x) = \binom{x+k-1}{x} q^x p^k \quad x = 0, 1, 2, \dots$$

رابطه‌ای مشابه رابطه بین متغیر تصادفی برنولی و دو جمله‌ای ما بین متغیر تصادفی هندسی نوع اول و متغیر تصادفی دو جمله‌ای منفی نوع اول موجود است. فرض کنید یک آزمایش را که متغیر تصادفی مربوط به آن از متغیر تصادفی هندسی نوع اول می‌باشد، k بار تکرار کنیم و نتایج حاصله را با یکدیگر جمع کنیم در این صورت تعداد شکستها قبل از رسیدن به k امین پیروزی را بدست می‌آوریم که همان متغیر تصادفی دو جمله‌ای منفی نوع اول باشد و $X_1, X_2, \dots, X_K$ متغیرهای تصادفی هندسی نوع اول باشند داریم:

$$Y = X_1 + X_2 + \dots + X_K$$

با استفاده از رابطه فوق امید ریاضی، واریانس و تابع مولد گشتاور متغیر تصادفی دو جمله‌ای منفی نوع اول را محاسبه می‌کنیم:

$$E[X] = E\left[\sum_{i=1}^k X_i\right] = \sum_{i=1}^k E[X_K] = \sum_{i=1}^k \frac{q}{p} = k \frac{q}{p}$$

$$\text{var}(X) = \text{var}\left(\sum_{i=1}^k X_i\right) = \sum_{i=1}^k \text{var}(X_i) = \sum_{i=1}^k \frac{q}{p^2} = k \frac{q}{p^2}$$

$$m_X(t) = E[e^{tX}] = E\left[e^{t\sum_{i=1}^k X_i}\right] = E[e^{tX_1} e^{tX_2} \dots e^{tX_K}]$$

$$E[e^{tX_1}] E[e^{tX_2}] \dots E[e^{tX_K}] =$$

$$= \left(\frac{p}{1-qe^t}\right) \left(\frac{p}{1-qe^t}\right) \dots \left(\frac{p}{1-qe^t}\right) = \left(\frac{p}{1-qe^t}\right)^k$$

۱۳-۶ مثال ۶: متغیر تصادفی X را تعداد دفعات به خطا رفتن گلوله قبل از اصابت ۳امین بار گلوله به هدف در مثال ۴ در نظر می‌گیریم، مطلوبست:

الف) تابع چگالی متغیر تصادفی X .

ب) احتمال اینکه پس از ۲ بار خطا رفتن گلوله برای سومین بار به هدف بزنیم.

ج) متوسط تعداد خطا رفتن گلوله قبل از اصابت سوم به هدف.

د) واریانس و تابع مولد گشتاور متغیر تصادفی X .

حل: الف) متغیر تصادفی X از تالچ چگالی دو جمله‌ای منفی نوع اول پیروی می‌کند بنابراین تابع چگالی آن با $p = \frac{2}{5}$ و $k = 3$ برابر است با:

$$f_X(x) = \binom{x+3-1}{x} \left(\frac{2}{5}\right)^x \left(\frac{3}{5}\right)^3$$

$$f(x=2) = \binom{4}{2} \left(\frac{2}{5}\right)^2 \left(\frac{3}{5}\right)^3 = 0.2 \quad \text{ب):}$$

$$E[X] = K \frac{q}{p} = 3 \frac{\frac{3}{5}}{\frac{2}{5}} = 4.5 \quad \text{ج):}$$

یعنی بطور متوسط پس از ۲ بار خطا رفتن گلوله باید بتوان هدف را مورد اصابت قرار داد.
همچنین:

$$\text{var}(X) = k \frac{q}{p^2} = 3 \frac{\frac{3}{5}}{\frac{4}{25}} = \frac{22.5}{4} = 5.625 \quad \text{د)}$$

$$m_X(t) = \left(\frac{p}{1-qe^t}\right)^k = \left(\frac{\frac{3}{5}}{1-\frac{2}{5}e^t}\right)^3 \quad \text{به همین ترتیب:}$$

۱۴-۶ ۵.۶ متغیر تصادفی دو جمله‌ای منفی (نوع دوم)

متغیر تصادفی X را برابر با تعداد آزمایشهای لازم برای رسیدن به k امین پیروزی در تکرار مستقل آزمایشهای برنولی در نظر می‌گیریم. در این صورت متغیر تصادفی X را دو جمله‌ای منفی نوع دوم می‌نامیم.

برای بدست آوردن تابع چگالی متغیر تصادفی دو جمله‌ای منفی نوع دوم ابتدا توجه کنید که تفاوت اصلی آن با متغیر تصادفی دو جمله‌ای نوع اول در این است که در نوع اول تعداد k بار پیروزی در محاسبه مقادیری که متغیر تصادفی قبول می‌کند در نظر گرفته نمی‌شود بنابراین اگر Y یک متغیر تصادفی دو جمله‌ای منفی نوع دوم باشد و X نوع اول رابطه زیر بین این دو متغیر تصادفی برقرار است:

$$Y = X + K \quad \begin{aligned} x &= 0, 1, 2, \dots \\ y &= k, k+1, k+2, \dots \end{aligned}$$

بنابراین تابع چگالی Y نیز با تغییر X به $Y - K$ در تالچ چگالی دو جمله‌ای منفی نوع اول بدست می‌آید:

$$f_X(x) = \binom{x+k-1}{x} q^x p^k \xrightarrow{x=y-k} f_Y(y) = \binom{y-1}{y-k} q^{y-k} p^k$$

$$\Rightarrow f_Y(y) = \binom{y-1}{k-1} q^{y-k} p^k \quad y = k, k+1, k+2, \dots$$

اگر $X_1, X_2, \dots, X_K$ مجموعه k متغیر تصادفی هندسی نوع دوم باشند و Y یک متغیر تصادفی دو جمله‌ای منفی نوع دوم باشد رابطه زیر بین آنها برقرار است:

$$Y = X_1 + X_2 + \dots + X_K$$

هر متغیر تصادفی X_i برابر است با تعداد آزمایشها برای رسیدن به اولین پیروزی، حال اگر نتایج حاصل از k آزمایش تصادفی هندسی نوع دوم را با یکدیگر جمع کنیم تعداد آزمایشهای لازم برای رسیدن به k امین پیروزی بدست می‌آید. که همان متغیر تصادفی دو جمله‌ای منفی می‌باشد. حال با توجه به رابطه فوق مقادیر امید ریاضی، واریانس و تابع مولد گشتاور را برای متغیر تصادفی دو جمله‌ای منفی نوع دوم بدست می‌آوریم:

$$E[X] = E\left[\sum_{i=1}^k X_k\right] = \sum_{i=1}^k E[X_k] = \sum_{i=1}^k \frac{1}{p} = \frac{k}{p}$$

$$\text{var}(X) = \text{var}\left(\sum_{i=1}^k X_k\right) = \sum_{i=1}^k \text{var}(X_k) = \sum_{i=1}^k \frac{q}{p^2} = K \frac{q}{p^2}$$

$$m_X(t) = E[e^{tX}] = E\left[e^{t\sum_{i=1}^k X_k}\right] = E[e^{tX_1} e^{tX_2} \dots e^{tX_K}]$$

$$= E[e^{tX_1}] E[e^{tX_2}] \dots E[e^{tX_K}]$$

$$= \left(\frac{pe^t}{1-qe^t}\right) \left(\frac{pe^t}{1-qe^t}\right) \dots \left(\frac{pe^t}{1-qe^t}\right) = \left(\frac{pe^t}{1-qe^t}\right)^k$$

مثال ۷: اگر در مثال ۴ Y را تعداد دفعات شلیک گلوله برای اصابت ۳ امین بار گلوله به هدف در نظر بگیریم مطلوبست:

الف) تابع چگالی متغیر تصادفی Y .

ب) احتمال اینکه پس از ۵ بار شلیک برای سومین بار هدف را بزنیم.

ج) متوسط تعداد شلیک برای اصابت سه بار گلوله به هدف.

د) واریانس و تابع مولد گشتاور متغیر تصادفی Y .

حل: الف) متغیر تصادفی Y از توزیع دو جمله‌ای منفی نوع دوم پیروی می‌ند بنابراین تابع چگالی آن بصورت زیر خواهد بود:

$$f_Y(y) = \binom{y-1}{2} \left(\frac{2}{5}\right)^{y-3} \left(\frac{3}{5}\right)^3 \quad y = 3, 4, 5, \dots$$

ب) می‌بایستی $f(y=5)$ را بدست بیاوریم:

$$f(y=5) = \binom{5-1}{2} \left(\frac{2}{5}\right)^{5-3} \left(\frac{3}{5}\right)^3 = 0.2$$

همانطور که ملاحظه می‌کنید مقدار $f(y=5)$ با مقدار $f(x=2)$ در مثال قبل برابر می‌باشد زیرا رابطه $Y = X + 3$ بین دو متغیر تصادفی X و Y در این دو مثال وجود دارد.

$$E[Y] = \frac{k}{p} = \frac{3}{\frac{1}{5}} = 5$$

ج):

یعنی بطور متوسط از هر ۵ بار شلیک به هدف می‌وان انتظار داشت که ۳ بار گلوله به هدف اصابت کند.

$$\text{var}(Y) = \frac{kq}{p^2} = 3 \frac{\frac{2}{5}}{\frac{1}{25}} = \frac{30}{9}$$

$$m_X(t) = \left(\frac{pe^t}{1-qe^t} \right)^k = \left(\frac{\frac{3}{5}e^t}{1-\frac{2}{5}e^t} \right)^3 \quad (د):$$

۱۷-۶ ۶.۶ متغیر تصافی فوق هندسی

فرض کنید ظرفی حاوی N توپ می‌باشد، که M تای آنها سفید می‌باشند. می‌خواهیم یک نمونه n تایی از این ظرف به تصادف و بدون جاگذاری خارج کنیم متغیر تصادفی X را تعداد توپهای سفیدی که در نمونه n تایی موجود می‌باشند در نظر می‌گیریم. در این صورت X یک متغیر تصادفی فوق هندسی نامیده می‌شود.

توجه کنید که در متغیر تصادفی فوق هندسی هر بار یک عدد توپ از ظرف خارج می‌کنیم که در اولین مرتبه، احتمال سفید بودن توپ $p = \frac{M}{N}$ می‌باشد، اما چون توپها را بدون جاگذاری خارج می‌کنیم در دومین انتخاب احتمال سفید بودن توپ به نتیجه آزمایش اول وابسته است به عبارت دقیقتر در انتخاب m توپ از ظرف، که هر انتخاب یک آزمایش برنولی است، بدلیل انتخاب توپها برون جاگذاری آزمایشها از یکدیگر مستقل نمی‌باشند.

اگر توپها را با جایگذاری انتخاب می‌کردیم در این صورت X یک متغیر تصادفی دو جمله‌ای با پارامترهای n ، $p = \frac{M}{N}$ می‌بوده.

برای بدست آوردن تابع چگالی متغیر تصادفی فوق هندسی فرض می‌نیم از n توپ انتخاب شده X تای آنها سفید باشند در این صورت X می‌تواند عددی بین 0 تا M باشد. انتخاب n توپ از N توپ به $\binom{N}{n}$ طریق ممکن است. همینطور سفید بودن X توپ به $\binom{M}{x}$ طریق و مابقی به $\binom{N-M}{n-x}$ طریق ممکن است در نتیجه احتمال $p(X=x)$ بصورت زیر بدست می‌آید:

$$f_X(x) = p(X=x) = \frac{\binom{M}{x} \binom{N-M}{n-x}}{\binom{N}{n}}$$

۱۷-۲ از $\binom{M}{x}$ نتیجه می‌شود که $0 \leq x \leq M$ و از $\binom{N-M}{n-x}$ نتیجه می‌شود که:

$$0 \leq n-x \leq N-M \Rightarrow n+M-N \leq x \leq n$$

نهایتاً بازده زیر برای مقادیری که متغیر تصادفی X قبول می‌کند بدست می‌آید:

$$\begin{cases} 0 \leq x \leq M \\ n+M-N \leq x \leq n \end{cases} \Rightarrow \max(0, n+M-N) \leq x \leq \min(M, n)$$

متغیر تصادفی فوق هندسی X را با نماد $X \sim HG(N, M, n)$ نمایش می‌دهیم.

امید ریاضی و واریانس متغیر تصادفی فوق هندسی برابر است با:

$$E[X] = n \frac{M}{N}$$

$$\text{var}(X) = n \frac{M}{N} \left(\frac{N-M}{N} \right) \left(\frac{N-n}{N-1} \right)$$

۱۸-۶ مثال ۸: در یک چاپخانه از هر ۵۰ برگ چاپ شده ۵ برگ بدلیل کیفیت نامناسب چاپ به دور ریخته می‌شود. برای کنترل کیفیت از هر ۵۰ برگ ۳ برگ کنترل می‌شود و اگر حداقل یک برگ کیفیت نامناسبی داشته باشند آنگاه سایر برگها نیز کنترل می‌شوند. احتمال اینکه مسوول کنترل کیفیت مجبور باشد هر ۵۰ برگ را کنترل کند چقدر می‌باشد؟

حل: متغیر تصادفی X را تعداد برگهای بدون کیفیت در یک نمونه ۳ تایی در نظر می‌گیریم. در این صورت X متغیر تصادفی فوق هندسی با پارامترهای $X \sim HG(50, 5, 3)$ می‌باشد.

$$p(X=x) = \frac{\binom{5}{x} \binom{45}{3-x}}{\binom{50}{3}} \quad 0 \leq x \leq 3$$

برای اینکه مجبور باشیم هر ۵۰ برگ را کنترل کنیم باید حداقل یک برگ انتخابی کیفیت نامناسبی داشته باشند که احتمال آن برابر است با: $p(x \geq 1)$

$$p(X \geq 1) = 1 - p(X=0) = 1 - \frac{\binom{5}{0} \binom{45}{3}}{\binom{50}{3}} = 1 - 0.723 = 0.276$$

متوسط تعداد برگهای با کیفیت نامناسب در نمونه عبارتست از:

$$E[X] = n \frac{M}{N} = 3 \frac{5}{50} = 0.3$$

$$\text{var}(X) = 3 \frac{5}{50} \left(\frac{50-5}{50} \right) \left(\frac{50-3}{50-1} \right) = 0.25$$

همچنین:

۱۹-۶ ۱.۶.۶ تقریب توزیع فوق هندسی بوسیله توزیع دو جمله‌ای

اگر در متغیر تصادفی فوق هندسی مقدار N در مقایسه با n بسیار بزرگ باشد، آنگاه دیگر انتخاب توپها بدون جاگذاری و با جاگذاری تقریباً معادل یکدیگر می‌باشند به همین دلیل به جای محاسبه تابع احتمال فوق هندسی می‌توانیم از تابع احتمال دو جمله‌ای استفاده کنیم که در این حالت از پارامترهای n ، $p = \frac{M}{N}$ استفاده می‌کنیم. در حالت کلی می‌بایستی شروط $N \rightarrow \infty$ ، $\frac{n}{N} \rightarrow 0$ برقرار باشند تا بتوان از توزیع دو جمله‌ای استفاده نمود.

مثال ۹: در یک شهر از میان ۱۰/۰۰۰ خانوار ۴۰۰۰ خانوار دارای فرزند پسر می‌باشند اگر یک نمونه ۵ تایی انتخاب کنیم احتمال اینکه دارای فرزند پسر نباشند چقدر است؟

حل: چون نسبت $\frac{5}{10.000}$ به صفر میل می‌کند و ۱۰۰۰۰ عدد بزرگی محسوب می‌شود می‌توان از تقریب توزیع فوق هندسی به دو جمله‌ای استفاده نمود:

$$p = \frac{M}{N} = \frac{4000}{10000} = \frac{2}{5}, \quad q = \frac{3}{5}, \quad n=5 \Rightarrow X \sim B\left(5, \frac{2}{5}\right)$$

$$p(X=x) = \binom{5}{x} \left(\frac{2}{5}\right)^x \left(\frac{3}{5}\right)^{5-x} \quad x=0, 1, 2, \dots, 5$$

پس:

$$= 0 \quad \text{سایر مقادیر}$$

احتمال اینکه فرزند پسری نداشته باشند برابر است با: $p(X=0)$.

$$p(X=0) = \binom{5}{0} \left(\frac{2}{5}\right)^0 \left(\frac{3}{5}\right)^5 = 0.07$$

۲۰-۶ ۷.۶ متغیر تصادفی پواسون

در بعضی از آزمایشها تعداد دفعات رخ دادن پیشامدی را می‌توان بدست آورد ولی تعداد دفعات عدم رخ دادن آنرا نمی‌توان بدست آورد مثلاً تعداد دفعاتی که در طول شبانه روز تلفن یک شرکت زنگ می‌زند شمرد ولی تعداد دفعاتی که تلفن زنگ نمی‌زند را نمی‌توان شمرد. همینطور تعداد خودروهای وارد شده به پمپ بنزین در مقایسه با خودروهای عبوری.

متغیر تصادفی X را به این صورت تعریف می‌کنیم:

$X =$ تعداد موفقیتها در یک فاصله پیوسته (زمان ، طول)

پارامتر تصادفی پواسون λ می‌باشد که برابر است با میانگین تعداد موفقیتها در طول یک بازه پیوسته. برای اینکه X متغیر تصادفی پواسون باشد می‌بایستی شرایط زیر برقرار باشد:

۱- در طول فاصله زمانی بسیار کوتاه احتمال رخ دادن یک موفقیت فقط متناسب با طول بازه زمانی باشد و بستگی به تعداد موفقیتها در خارج از فاصله زمانی نداشته باشد.

۲- احتمال روی دادن بیش از یک موفقیت در یک فاصله زمانی کوتاه تقریباً برابر صفر باشد.

۳- تعداد موفقیتهایی که در یک فاصله زمانی مشخص روی می‌دهد از تعداد موفقیتهایی که در یک فاصله زمانی دیگر رخ می‌دهد مستقل باشد. توجه کنید که در شروط بالا منظور از فاصله زمانی هر بازه پیوسته مثل طول یا یک ناحیه مشخص می‌باشد.

تابع احتمال متغیر تصادفی پواسون با پارامتر λ بصورت زیر می‌باشد:

$$f_X(x) = \frac{e^{-\lambda} \lambda^x}{x!} \quad x = 0, 1, 2, \dots$$

$$= 0 \quad \text{سایر مقادیر}$$

مقادیر امید ریاضی، واریانس و تابع مولد گشتاور عبارتند از:

$$E[X] = \lambda$$

$$\text{var}(X) = \lambda$$

$$m_X(t) = e^{-\lambda(1-e^t)}$$

معمولاً پارامتر λ به صورت $\lambda = rt$ در نظر گرفته می‌شود که در آن t طول بازه مورد نظر و r نرخ وقوع پیشامد در بازه مربوطه است.

مثال ۱۰: یک تایپیست به طور متوسط در هر ۲ صفحه ۵ غلط تایپی دارد مطلوبست:

الف) احتمال اینکه در یک صفحه اصلاً غلط تایپی نداشته باشد؟

ب) احتمال اینکه حداقل ۳ غلط تایپی در ۲ صفحه داشته باشد؟

حل: الف) در یک صفحه ۲/۵ غلط به طور متوسط موجود می‌باشد بنابراین $\lambda = 2/5 \times 1$ و داریم:

$$f_X(x) = \frac{(2/5)^x e^{-2/5}}{x!}$$

برای اینکه اصلاً غلط تایپی نداشته باشیم بایستی $f(x=0)$ را بدست بیاوریم:

$$f(x=0) = \frac{(2/5)^0 e^{-2/5}}{0!} = e^{-2/5} = 0.6703$$

ب) در این حالت $\lambda = 5$ می‌باشد و داریم:

$$f_X(x) = \frac{5^x e^{-5}}{x!}$$

$$p(X \geq 3) = 1 - (p(X=0) + p(X=1) + p(X=2))$$

$$= 1 - (0.0067 + 0.0337 + 0.0842) = 0.8754$$

مد توزیع پواسون همانند توزیع دو جمله‌ای بدست می‌آید یعنی:

$$\begin{cases} f(x_M) \geq f(x_M - 1) & (1) \\ f(x_M) \geq f(x_M + 1) & (2) \end{cases}$$

از (۱) نتیجه می‌شود: قرار دادن $x_M = \hat{x}$

$$\frac{\lambda^{\hat{x}} e^{-\lambda}}{\hat{x}!} \geq \frac{\lambda^{\hat{x}-1} e^{-\lambda}}{(\hat{x}-1)!} \Rightarrow x \leq \lambda$$

$$\lambda - 1 \leq x \leq \lambda$$

به همین نسبت از رابطه (۲) بدست می‌آید $x \leq \lambda - 1$ در نتیجه:

حال اگر λ عددی صحیح باشد λ ، $\lambda - 1$ هر دو مد می‌باشند و اگر λ عددی صحیح نباشد بین λ و $\lambda - 1$ عددی صحیح موجود است که مد می‌باشد.

مثال ۱۱: برای مثال قبل مقدار مد را برای بند الف و ب محاسبه کنید؟

الف) $\lambda = 2/5$ ، $\lambda - 1 = 1/5$ بنابراین مد برابر است با $x_M = 2$.

ب) $\lambda = 5$ ، $\lambda - 1 = 4$ بنابراین مد برابر است با $x_M = 5$.

۶.۷.۲ تقریب توزیع دو جمله‌ای بوسیله توزیع پواسون

فرض کنید در توزیع دو جمله‌ای $n \rightarrow \infty$ ، $p \rightarrow 0$ می‌توان به جای محاسبه احتمال با استفاده از تابع احتمال دو جمله‌ای از تابع احتمال پواسون استفاده نمود. در این حالت توزیع دو جمله‌ای را با توزیع پواسون با پارامتر $\lambda = np$ تقریب می‌زنیم.

توجه کنید که برای بدست آوردن مقدار λ امید ریاضی تابع پواسون را با امید ریاضی تابع احتمال دو جمله‌ای برابر قرار می‌دهیم که در نتیجه $\lambda = np$ خواهد شد.

مثال ۱۲: در طول یک روز تعداد ۱۰۰/۰۰۰ خودرو از جلوی پمپ بنزین عبور می‌کند که تعداد ۱۰۰ خودرو وارد پمپ بنزین می‌شوند و مطلوبست

احتمال اینکه در طول روز حداقل ۵۰ خودرو وارد پمپ بنزین شوند؟

حل: با محاسبه $p = \frac{100}{100/000} = 0/001$ که عددی بسیار کوچک می‌باشد و چون n مقداری بزرگ می‌باشد می‌توان از تقریب توزیع دو

جمله‌ای به پواسون استفاده نمود:

$$\lambda = np = 100/000 \times 0/001 = 100$$

$$f_X(x) = \frac{100^x e^{-100}}{x!}$$

$$p(X \geq 50) = 1 - p(X < 50) \sim 0/999$$

۶.۶.۸ توزیع یکنواخت

متغیر تصادفی X دارای توزیع یکنواخت می‌باشد اگر احتمال رخ دادن هر یک از نقاط $x=1, x=2, \dots, x=n$ با یکدیگر برابر باشد. تابع چگالی متغیر تصادفی یکنواخت برابر است با:

$$f_X(x) = \frac{1}{n} \quad x=1, 2, \dots, n$$

سایر مقادیر = ۰

امید، واریانس و تابع گشتاور توزیع یکنواخت عبارتند از:

$$E[X] = \sum_{i=1}^n x \frac{1}{n} = \frac{1}{n} \left(\frac{n(n+1)}{2} \right) = \frac{n+1}{2}$$

$$E[X^2] = \sum_{i=1}^n x_i^2 \frac{1}{n} = \frac{1}{n} \left(\frac{n(n+1)(2n+1)}{6} \right) = \frac{(n+1)(2n+1)}{6}$$

به همین ترتیب:

پس:

$$\text{var}(X) = E[X^2] - E^2[X] = \frac{(n+1)(2n+1)}{6} - \left(\frac{n+1}{2}\right)^2 = \frac{n^2-1}{12}$$

$$m_X(t) = E[e^{tx}] = \sum_{x=1}^n \frac{e^{tx}}{n} = \frac{1}{n} \sum_{i=1}^n (e^t)^x = \frac{1}{n} e^t \left(\frac{1-e^{nt}}{1-e^t} \right)$$

مثال ۱۳: یک تاس را پرتاب می‌کنیم X را متغیر تصادفی نتیجه حاصل از پرتاب در نظر می‌گیریم مطلوبست:

الف) تابع چگالی احتمال X .

ب) احتمال آمدن عدد زوج.

ج) امید، واریانس و تابع مولد گشتاور X .

حل: الف) X یک متغیر تصادفی یکنواخت است. بنابراین:

$$f(x) = \frac{1}{6} \quad x=1, 2, \dots, 6$$

= ۰ سایر مقادیر

ب):

$$p(X = \text{عدد زوج}) = p(X=2) + p(X=4) + p(X=6) = \frac{3}{6} = \frac{1}{2}$$

$$E[X] = \frac{6+1}{2} = 3.5$$

ج):

$$\text{var}(X) = \frac{36-1}{12} = 2.91$$

$$m_X(t) = \frac{e^t(1-e^{6t})}{6(1-e^t)}$$

فصل هفتم

در فصل قبل برخی توزیع‌های معروف و کاربردی را برای متغیر تصادفی گسسته ارایه نمودیم و خواص و ویژگیهای هر یک را بررسی نمودیم. در این فصل به بررسی توزیعهای مبتنی بر متغیر تصادفی پیوسته می‌پردازیم.

۱.۷ توزیع یکنواخت پیوسته

ساده‌ترین توزیع پیوسته، توزیع یکنواخت می‌باشد. فرض کنید تمام نقاط در بازه (a, b) دارای امکان وقوع یکسان باشند، در این صورت متغیر تصادفی X را که برد آن مقادیر موجود در بازه (a, b) می‌باشد، متغیر تصادفی یکنواخت می‌نامند. که با نماد $X \sim U(a, b)$ نشان داده می‌شود. تابع چگالی احتمال متغیر تصادفی یکنواخت برابر است با:

$$f_X(x) = \begin{cases} \frac{1}{b-a} & a < x < b \\ 0 & \text{سایر مقادیر} \end{cases}$$

و به این ترتیب بدست می‌آید که چون X یکنواخت می‌باشد در نظر می‌گیریم:

$$f_X(x) = c \quad a < x < b$$

در نتیجه:

$$\int_{-\infty}^{+\infty} f_X(x) dx = 1 \Rightarrow \int_a^b c dx = c x \Big|_a^b = c(b-a) = 1 \Rightarrow c = \frac{1}{b-a}$$

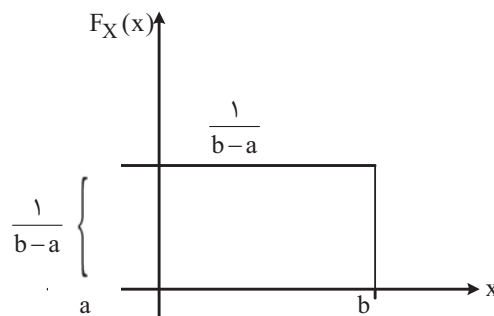
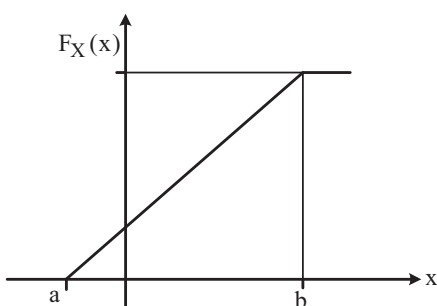
تابع توزیع متغیر تصادفی یکنواخت X برابر است با:

$$f_X(x) = \int_{-\infty}^t f_X(x) dx = \int_a^t \frac{1}{b-a} dx = \frac{t-a}{b-a}$$

$$f_X(x) = \begin{cases} 0 & t < a \\ \frac{t-a}{b-a} & a \leq t \leq b \\ 1 & t > b \end{cases}$$

پس:

۷-۲ در شکل زیر نمودار تابع چگالی و تابع توزیع متغیر تصادفی یکنواخت نشان داده شده است:



مقادیر امید ریاضی و واریانس و تابع مولد گشتاور را برای متغیر تصادفی یکنواخت X محاسبه می‌کنیم:

$$E[X] = \int_a^b x \frac{1}{b-a} dx = \frac{1}{b-a} \frac{x^2}{2} \Big|_a^b =$$

$$= \frac{1}{b-a} \left(\frac{b^2}{2} - \frac{a^2}{2} \right) = \frac{(b-a)(a+b)}{2(b-a)} = \frac{a+b}{2}$$

$$E[X^2] = \int_a^b x^2 \frac{1}{b-a} dx = \frac{1}{b-a} \left(\frac{x^3}{3} \right) \Big|_a^b$$

$$= \frac{(b-a)(b^3 - ab + a^3)}{3(b-a)} = \frac{a^3 + ab + b^3}{3}$$

پس:

$$\text{var}(X) = E[X^2] - E^2[X] = \frac{a^3 + ab + b^3}{3} - \frac{a^2 + 2ab + b^2}{4} = \frac{(b-a)^2}{12}$$

$$m_X(t) = E[e^{tX}] = \int_a^b e^{tX} \frac{1}{b-a} dx = \frac{1}{b-a} \left(\frac{e^{tX}}{t} \right) \Big|_a^b = \frac{e^{tb} - e^{ta}}{t(b-a)}$$

با فرض $a \leq c < d \leq b$ مقدار احتمال $c \leq x \leq d$ برابر است با:

$$p(c \leq x \leq d) = \int_c^d \frac{1}{b-a} dx = \frac{d-c}{b-a}$$

ملاحظه می‌کنید که مطابق با تعریف متغیر تصادفی یکنواخت مقدار احتمال تنها به طول بازه (c, d) وابسته است و نه به مقادیر c و d .

۷-۳ مثال ۱: نمرات دانشجویان یک کلاس در درس آمار به طور یکنواخت در فاصله ۱۲ الی ۲۰ توزیع شده است مطلوبست:

الف) تابع توزیع احتمال:

ب) احتمال قرارگیری نمرات در بازه $(18-20)$ ، $(13-15)$.

ج) میانگین نمرات کلاس، واریانس.

حل: الف) متغیر تصادفی X دارای توزیع یکنواخت است بنابراین:

$$f_X(x) = \begin{cases} \frac{1}{20-12} = \frac{1}{8} & 12 \leq X \leq 20 \\ 0 & \text{سایر مقادیر} \end{cases}$$

ب):

$$P(18 \leq X \leq 20) = \int_{18}^{20} \frac{1}{8} dx = \frac{20-18}{8} = \frac{2}{8} = \frac{1}{4}$$

$$P(13 \leq X \leq 15) = \int_{13}^{15} \frac{1}{8} dx = \frac{15-13}{8} = \frac{2}{8} = \frac{1}{4}$$

ج):

$$E[X] = \frac{20+12}{2} = 16$$

$$\text{var}(X) = \frac{(20-12)^2}{12} = \frac{64}{12} = \frac{16}{3}$$

۷-۴ ۲.۶ متغیر تصادفی نمایی

یک فرایند پواسون با پارامتر λ که برای یک واحد زمان نظاره می‌شود، را در نظر بگیرید اگر زمان شروع فرایند صفر $(t=0)$ باشد و T مدت زمانی باشد که باید بگذرد تا اولین پیشامد رخ دهد در این صورت T یک متغیر تصادفی نمایی با پارامتر λ نامیده می‌شود.

تابع توزیع نمایی را با استفاده از تعریف بدست می‌آوریم:

می‌دانیم $f_X(t) = p(X \leq t)$ حال اگر $t < 0$ باشد بنا به تعریف زمان منفی معنا ندارد. بنابراین اگر $t < 0$ داریم: $p(X \leq t) = 0$ یعنی:

$$f_X(t) = 0 \quad t < 0$$

حال فرض می‌کنیم $t \geq 0$ باشد، احتمال اینکه هیچ پیشامدی در بازه $(0, t)$ رخ ندهد بنابه تابع چگالی پواسون برابر است با:

$$f_X(x=0) = \frac{(\lambda t)^x e^{-\lambda t}}{x!} \Big|_{x=0} = e^{-\lambda t}$$

بنابراین $p(X > t) = e^{-\lambda t}$ اما $p(X \leq t) = 1 - p(X > t)$ و داریم:

$$F_X(t) = p(X \leq t) = 1 - p(X > t) = 1 - e^{-\lambda t}$$

در نتیجه تابه توزیع متغیر تصادفی نمایی بصورت زیر بدست می‌آید:

$$F_X(t) = \begin{cases} 1 - e^{-\lambda t} & t > 0 \\ 0 & t < 0 \end{cases}$$

با مشتق‌گیری از $F_X(t)$ تابع چگالی را بدست می‌آوریم:

$$f_X(x) = \begin{cases} \lambda e^{-\lambda x} & x > 0 \\ 0 & x \leq 0 \end{cases}$$

۷-۵ برای بدست آوردن مقادیر امید و واریانس از تابع مولد گشتاور استفاده می‌کنیم:

$$m_X(t) = E[e^{tx}] = \int_0^{\infty} e^{tx} \lambda e^{-\lambda x} dx = \lambda \int_0^{\infty} e^{tx} \lambda e^{-x(\lambda-t)} dx$$

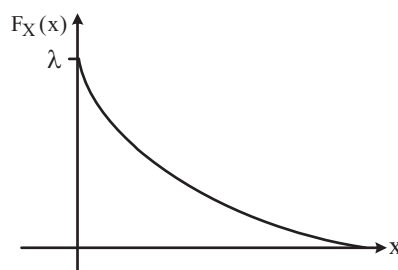
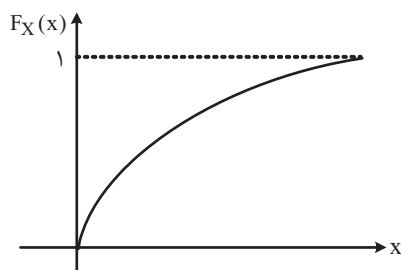
$$= \lambda \left(\frac{1}{t-\lambda} e^{-x(\lambda-t)} \right) \Big|_0^{\infty} = \frac{-\lambda}{t-\lambda} = \frac{\lambda}{\lambda-t} \quad (t < \lambda)$$

$$E[X] = m'_X(t=0) = \frac{\lambda}{(\lambda-t)^2} \Big|_{t=0} = \frac{1}{\lambda}$$

$$E[X^2] = m''_X(t=0) = \frac{2\lambda}{(\lambda-t)^3} \Big|_{t=0} = \frac{2}{\lambda^2}$$

$$\text{var}(X) = E[X^2] - E^2[X] = \frac{2}{\lambda^2} - \frac{1}{\lambda^2} = \frac{1}{\lambda^2}$$

شکل زیر نمودار تابع توزیع و چگالی را برای متغیر تصادفی نمایی می‌دهد:



متغیر تصادفی نمایی با نماد $X \sim E(\lambda)$ نمایش داده می‌شود و معمولاً از آن به عنوان مدلی برای عمر قطعات و سیستم‌ها استفاده می‌شود.

۷-۶ مثال ۲: اگر لامپهای تولید شده توسط یک کارخانه به طور متوسط هر ۶ ماه یکبار بسوزند مطلوبست:

(الف) تابع چگالی احتمال برای مدت زمان کارکرد لامپها.

(ب) احتمال اینکه از زمان خرید، یک لامپ حداقل ۲ ماه کار کند؟

(ج) احتمال اینکه یک لامپ قبل از ۴ ماه بسوزد؟

(د) اگر ۱۰ عدد لامپ خریداری کنیم احتمال اینکه حداقل ۲ عدد از لامپها قبل از ۴ ماه بسوزند؟

ه) میانگین طول عمر، واریانس و تابع مولد گشتاور را برای متغیر تصادفی X که مدت زمان کارکرد لامپها می باشد بدست بیاورید.

۷-۷ حل: الف) متغیر تصادفی X را طول عمر لامپهای خریداری شده در نظر می گیریم در این صورت X یک متغیر تصادفی نمایی با پارامتر λ می باشد. برای بدست آوردن پارامتر λ ابتدا هر واحد زمانی را یک ماه در نظر می گیریم، λ برابر است با تعداد پیشامدها در یک واحد زمانی که در

اینجا چون هر ۶ ماه یک لامپ می سوزد معادل است با اینکه $\frac{1}{6}$ لامپ در هر یک ماه می سوزد یعنی $\lambda = \frac{1}{6}$ و تابع چگالی برابر است با:

$$f_X(x) = \frac{1}{6} e^{-\frac{1}{6}x} \quad x > 0$$

سایر مقادیر = 0

ب) احتمال اینکه حداقل ۲ ماه طول بکشد تا یک لامپ بسوزد برابر است با:

$$p(X > 2) = \int_2^{\infty} \frac{1}{6} e^{-\frac{x}{6}} dx = \frac{1}{6} (-6 e^{-\frac{x}{6}}) \Big|_2^{\infty}$$

$$= \frac{1}{6} (0 - (-6 e^{-\frac{2}{6}})) = e^{-\frac{1}{3}} = 0.716$$

ج) احتمال اینکه یک لامپ قبل از ۴ ماه بسوزد برابر است با:

$$p(X \leq 4) = \int_0^4 \frac{1}{6} e^{-\frac{x}{6}} dx = \frac{1}{6} (-6 e^{-\frac{x}{6}}) \Big|_0^4 = 1 - e^{-\frac{4}{6}} = 0.486$$

۷-۸ د) متغیر Y را بصورت زیر تعریف می کنیم:

$Y =$ تعداد لامپها از بین ۱۰ لامپ که در کمتر از ۴ ماه می سوزند.

طبق تعریف Y یک متغیر تصادفی دو جمله ای با پارامتر $n = 10$ و $p = 0.486$ می باشد. پارامتر p از بند (ج) بدست می آید. احتمال اینکه هر لامپ در کمتر از ۴ ماه بسوزد برابر 0.486 می باشد، بنابراین می بایستی پارامتر p برابر 0.486 انتخاب شود. حال با احتمال اینکه حداقل ۲ عدد از لامپها قبل از ۴ ماه بسوزند برابر است با:

$$p(Y \geq 2) = 1 - p(Y < 2) = 1 - [p(Y=1) + p(Y=0)]$$

$$= 1 - \left[\binom{10}{1} (0.486)^1 (0.514)^9 + \binom{10}{0} (0.486)^0 (0.514)^{10} \right]$$

$$= 1 - (0.012 + 0.0012) = 0.988$$

$$f_X(x) = \binom{10}{y} (0.486)^y (0.514)^{10-y}$$

پس:

اگر امید ریاضی $f_Y(y)$ را محاسبه کنیم داریم: $E[X] = np = 10 \times 0.486 = 4.86$ یعنی تقریباً از هر ۱۰ عدد لامپ خریداری شده بطور متوسط ۵ لامپ در کمتر از ۴ ماه می سوزند.

ه) میانگین طول عمر لامپها عبارتست از: $E[X] = \frac{1}{\lambda} = \frac{1}{\frac{1}{6}} = 6$ که با توجه به اینکه واحد زمان را هر یک ماه در نظر گرفتیم بنابراین میانگین طول عمر هر لامپ ۶ ماه می باشد که از صورت مثال نیز همین انتظار می رفت.

$$\text{var}(X) = \frac{1}{\lambda^2} = \frac{1}{\left(\frac{1}{6}\right)^2} = 36$$

$$m_X(t) = \frac{\lambda}{\lambda - t} = \frac{\frac{1}{6}}{\frac{1}{6} - t} = \frac{1}{1 - 6t}$$

۹-۶.۲.۱ رابطه توزیع هندسی و توزیع نمایی

توزیع هندسی طبق تعریف عبارتست از تعداد آزمایشها قبل از رسیدن به اولین پیروزی توزیع نمایی نیز از جهت تعریف تشابه زیادی با توزیع هندسی دارد. توزیع نمایی هم عبارتست از مدت زمانی که در یک فرایند پواسون با پارامتر λ باید سپری شود تا اولین پیشامد رخ دهد توجه کنید که میانگین توزیع هندسی $\frac{1}{p}$ و توزیع نمایی $\frac{1}{\lambda}$ می باشد. متغیرهای تصادفی هندسی و نمایی در یک خاصیت ویژه مشترک می باشند که هیچ متغیر تصادفی گسسته یا پیوسته دیگری این حالت را ندارد. ویژگی فوق به بی حافظگی معروف است و به صورت زیر می باشد:

برای توزیع نمایی: $p(X > a + b | X > a) = p(X > b)$

برای توزیع هندسی: $p(X = a + b | X \geq a) = p(X = b)$

اثبات برای توزیع نمایی: فرض می نیم A برابر با پیشامد $X > a$ و B برابر پیشامد $X > b$ و C برابر پیشامد $X > a + b$ باشد در این صورت

$$p(C|A) = \frac{p(C \cap A)}{p(A)} \quad \text{باید ثابت کنیم: } p(C|A) = p(B) \text{ می دانیم:}$$

از آنجا که پیشامد C زیر مجموعه ای از پیشامد A است $(\{X > a + b\} \subset \{X > a\})$ بنابراین $C \cap A = C$ و داریم:

$$p(C|A) = \frac{p(C \cap A)}{p(A)} = \frac{p(C)}{p(A)}$$

۱۰-۷ حال مقادیر $p(A)$ و $p(B)$ و $p(C)$ را بدست می آوریم:

$$p(A) = p(X > a) = \int_a^{\infty} \lambda e^{-\lambda x} dx = e^{-\lambda a}$$

$$p(B) = p(X > b) = \int_b^{\infty} \lambda e^{-\lambda x} dx = e^{-\lambda b}$$

$$p(C) = p(X > a + b) = \int_{a+b}^{\infty} \lambda e^{-\lambda x} dx = e^{-\lambda(a+b)}$$

$$p(C|A) = \frac{p(C)}{p(A)} = \frac{e^{-\lambda(a+b)}}{e^{-\lambda a}} = e^{-\lambda b} = p(B)$$

بی حافظگی به این مفهوم است که اگر قطعه ای یا سیستمی به مدت a واحد زمان کار کرده باشد، احتمال اینکه حداقل b واحد زمان دیگر نیز کار کند $(a + b)$ برابر است با احتمال اینکه قطعه یا سیستم از لحظه صفر بخواهد حداقل b واحد زمان کار کند.

۱۱-۷ مثال ۳: در یک کارخانه ماشین های تولیدی هر یک ماه نیازمند سرویس تعمیرات باشند. اگر یک ماشین تولیدی ۶ ماه بدون تعمیرات کار

کرده باشد احتمال اینکه در طول ماه بعد نیازمند تعمیرات باشد چقدر است؟

حل: متغیر تصادفی X را مدت زمان لازم قبل از اولین تعمیر در نظر می گیریم در این صورت X یک متغیر تصادفی نمایی با پارامتر $\lambda = 1$ می باشد. می بایستی مقدار احتمال زیر را محاسبه کنیم:

$$p(X > 6 + 1 | X > 6) = p(X > 7 | X > 6)$$

$$= p(X > 1) = \int_1^{+\infty} e^{-x} dx = e^{-1} = 0.36$$

۱۲-۷ ۳.۶ توزیع گاما

اگر متغیرهای تصادفی $X_1, X_2, \dots, X_n$ دارای توزیع نمایی با پارامتر λ باشند در این صورت متغیر تصادفی Y را که برابر است با مجموع متغیرهای تصادفی X_1 تا X_n ، یک متغیر تصادفی گاما می‌نامیم.

$$Y = X_1 + X_2 + \dots + X_n$$

تابع چگالی متغیر تصادفی گاما برابر است با:

$$f_X(x) = \frac{(\lambda x)^{n-1}}{\Gamma(n)} \lambda e^{-\lambda x} \quad x \geq 0$$

سایر مقادیر $= 0$

که در آن $T(n)$ تابع گاما می‌باشد و به صورت زیر تعریف می‌شود:

$$\Gamma(n) = \int_0^{\infty} x^{n-1} e^{-x} dx$$

$\Gamma(n)$ به ازای هر $n > 0$ موجود است و روابط زیر برای تابع گاما برقرار است:

$$1- \Gamma(n+1) = n \Gamma(n)$$

$$2- \Gamma(n) = (n-1)! \text{ اگر } n \text{ عدد طبیعی باشد}$$

$$3- \Gamma\left(\frac{1}{2}\right) = \sqrt{\pi}$$

از آنجا که متغیر تصادفی گاما مجموع n متغیر تصادفی نمایی است پس می‌توان مقادیر امید، واریانس و تابع مولد گشتاور را با استفاده از این خاصیت بدست آورد:

$$E[X] = E\left[\sum_{i=1}^n X_i\right] = \sum_{i=1}^n E[X_i] = \sum_{i=1}^n \frac{1}{\lambda} = \frac{n}{\lambda}$$

$$\text{var}(X) = \text{var}\left(\sum_{i=1}^n X_i\right) = \sum_{i=1}^n \text{var}(X_i) = \sum_{i=1}^n \frac{1}{\lambda^2} = \frac{n}{\lambda^2}$$

$$m_X(t) = E[e^{tX}] = E\left[e^{t \sum_{i=1}^n X_i}\right] = E[e^{tX_1}] E[e^{tX_2}] \dots E[e^{tX_n}]$$

$$= \left(\frac{\lambda}{\lambda-t}\right) \left(\frac{\lambda}{\lambda-t}\right) \dots \left(\frac{\lambda}{\lambda-t}\right) = \left(\frac{\lambda}{\lambda-t}\right)^n$$

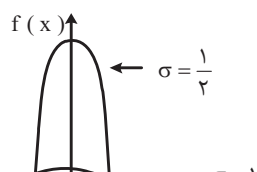
۷-۱۳ توزیع نرمال

مهمترین توزیع پیوسته که آنرا بررسی می‌کنیم توزیع نرمال می‌باشد، بسیاری از پدیده‌های طبیعی مثل قد و وزن افراد، نمرات درسی، میزان محصول در طول سال از توزیع نرمال پیروی می‌کنند.

متغیر تصادفی X دارای توزیع نرمال است و آنرا بصورت $X \sim N(\mu, \sigma^2)$ نمایش می‌دهیم اگر تابع چگالی احتمال بفرم زیر باشد:

$$f_X(x) = \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} \quad -\infty < x < +\infty, \mu \in \mathbb{R}, \sigma > 0$$

همانطور که ملاحظه می‌نید دو پارامتر توزیع نرمال همان میانگین و واریانس متغیر تصادفی X می‌باشند. مد، میانه و میانگین توزیع نرمال با یکدیگر برابر می‌باشند. در شکل زیر توزیع نرمال با مقادیر واریانس متفاوت نشان داده شده است.



تغیر میانگین در توزیع نرمال تنها نمودار تابع را به سمت راست یا چپ منتقل می‌کند. منحنی نرمال دارای خواص زیر می‌باشد:

۱- نقاط عطف منحنی عبارتند از $x_1 = \mu + \sigma$, $x_2 = \mu - \sigma$.

۲- منحنی نسبت به خط $x = \mu$ متقارن است بنابراین:

$$f_X(\mu - a) = f_X(\mu + a) \quad (a > 0)$$

$$\begin{cases} p(X > a) = p(X < -a) \\ 1 - F_X(x) = F_X(-x) \end{cases}$$

$$\int_{-\infty}^{+\infty} f(x) dx = 1 \quad \text{۳- مساحت سطح زیر منحنی نمودار نرمال برابر واحد می‌باشد. زیرا بوضوح داریم:}$$

۱۴-۷ برای بدست آوردن میانگین و واریانس از تابع مولد گشتاور استفاده می‌کنیم که بصورت زیر بدست می‌آید:

$$m_X(t) = E[e^{tX}] = \int_{-\infty}^{+\infty} e^{tx} \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} dx$$

$$= \frac{1}{\sigma \sqrt{2\pi}} \int_{-\infty}^{+\infty} e^{tx - \frac{(x-\mu)^2}{2\sigma^2}} dx$$

با تبدیل توان e به صورت مربع کامل داریم:

$$tx - \frac{(x-\mu)^2}{2\sigma^2} = \frac{1}{2\sigma^2} [x^2 - 2\mu x + \mu^2 - 2\sigma^2 t x + x^2]$$

$$= \frac{1}{2\sigma^2} [(x - \mu - \sigma^2 t)^2 - 2\sigma^2 t \mu - \sigma^4 t^2]$$

$$= \frac{1}{2\sigma^2} [(x - \mu - \sigma^2 t)^2] + t\mu + \frac{\sigma^2 t^2}{2}$$

به این ترتیب:

$$m_X(t) = e^{\mu t + \frac{(\sigma^2 t^2)}{2}} \underbrace{\int_{-\infty}^{+\infty} \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x-\mu-\sigma^2 t)^2}{2\sigma^2}} dx}_{g_X(x)}$$

مقدار انتگرال $\int_{-\infty}^{+\infty} g(x) dx$ برابر یک می‌باشد زیرا $g_X(x)$ خود یک تابع نرمال با پارامترهای $(\mu + \sigma^2 t, \sigma^2)$ می‌باشد. بنابراین مقدار

تابع مولد گشتاور برابر است با:

$$m_X(t) = e^{\mu t + \frac{\sigma^2 t^2}{2}}$$

$$E[X] = m'_X(t) \Big|_{t=0} = \left(\mu + \frac{\sigma^2}{2} (2t) \right) e^{\mu t + \frac{\sigma^2 t^2}{2}} \Big|_{t=0} = \mu \quad \text{حال داریم:}$$

$$E[X^2] = m_X''(t) \Big|_{t=0} = \sigma^2 (e^{\mu t + \frac{\sigma^2 t^2}{2}}) \left(\mu + \frac{\sigma^2}{2} (2t) \right) e^{\mu t + \frac{\sigma^2 t^2}{2}} \Big|_{t=0} = \sigma^2 + \mu^2$$

$$\Rightarrow \text{var}(X) = E[X^2] - E^2[X] = \sigma^2 + \mu^2 - \mu^2 = \sigma^2$$

۷-۱۵ تابع توزیع متغیر تصادفی نرمال X عبارتست از:

$$F_X(t) = \int_{-\infty}^t \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}} dx$$

برای انتگرال فوق یک تابع ائلیه نمی توان بدست آورد به همین دلیل مجبور هستیم از روشهای عددی یک جدول برای مقادیر متفاوت t بدست بیاوریم. اما از آنجا که دو پارامتر μ و σ^2 نیز متغیر می باشند می بایستی روشی بدست بیاوریم که بتوان مقدار احتمال را بدون وابستگی به μ و σ^2 بدست آورد. متغیر تصادفی Z را بصورت زیر تعریف می کنیم:

$$Z = \frac{X - \mu}{\sigma}$$

بوضوح Z متغیر تصادفی نرمال با میانگین ۰ و واریانس ۱ می باشد. $Z \sim N(0, 1)$ متغیر تصادفی Z را نرمال استاندارد می نامند. (که در فصل اول نیز برای نمونه های گرفته شده از یک جامعه معرفی شد).

$$E[Z] = E\left[\frac{X - \mu}{\sigma}\right] = \frac{1}{\sigma} [E[X] - \mu] = \frac{1}{\sigma} (\mu - \mu) = 0$$

$$\text{var}(Z) = \text{var}\left(\frac{X - \mu}{\sigma}\right) = \frac{1}{\sigma^2} \text{var}(X - \mu) = \frac{\sigma^2}{\sigma^2} = 1$$

$$m_Z(t) = e^{\frac{t^2}{2}}$$

حال با استفاده از تغییر متغیر $Z = \frac{X - \mu}{\sigma}$ می توان به سادگی مقادیر احتمال را از روی جدول توزیع متغیر تصادفی نرمال استاندارد بدست آورد.

تابع چگالی متغیر تصادفی نرمال استاندارد که آنرا با $\eta_Z(z)$ یا $\phi(z)$ نمایش می دهیم برابر است با:

$$\eta_Z(z) = \phi(z) = \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}} \quad -\infty < Z < +\infty$$

تابع توزیع متغیر تصادفی نرمال استاندارد برابر است با:

$$\phi(Z) = N_Z(z) = p_Z(Z \leq z) = \int_{-\infty}^z \frac{1}{\sqrt{2\pi}} e^{-\frac{z^2}{2}} dx$$

جدول ضمیمه مقادیر تابع توزیع را برای $-\frac{3}{5} \leq Z \leq \frac{3}{5}$ نشان می دهد.

۷-۱۶ مثال ۴: با استفاده از جدول مقادیر احتمالات زیر را بدست بیاورید؟

(الف) $p(Z < 0)$, $p(Z > 0)$

(ب) $p(Z < 1)$, $p(|Z| < \frac{3}{5})$

(ج) $p(-1 < Z < 0)$, $p(1 < Z < 3)$

حل: الف)

$$p(Z < 0) = \Phi(0) = \frac{1}{2}$$

$$p(Z > 0) = 1 - \Phi(0) = 1 - \frac{1}{2} = \frac{1}{2}$$

(ب)

$$p(Z < 1) = \Phi(1) = 0.8413$$

$$p(|Z| < \frac{3}{2}) = p(-\frac{3}{2} < Z < \frac{3}{2}) = \Phi(\frac{3}{2}) - \Phi(-\frac{3}{2})$$

$$= 0.9332 - 0.0668 = 0.8664$$

(ج)

$$p(-1 < Z < 0) = \Phi(0) - \Phi(1) = \frac{1}{2} - 0.8413 = 0.1587$$

$$p(1 < Z < 3) = \Phi(3) - \Phi(1) = 0.9987 - 0.8413 = 0.1574$$

با توجه به مثال فوق می‌وان خواص زیر را برای متغیر تصادفی نرمال استاندارد به دست آورد:

$$\Phi(-\infty) = 0 \quad 1.$$

$$\Phi(+\infty) = 1 \quad 2.$$

$$\Phi(a) = 1 - \Phi(-a) \quad 3.$$

۱۷-۷.۴.۶ محاسبه مقادیر احتمال متغیر تصادفی نرمال

فرض کنید X یک متغیر تصادفی نرمال با میانگین μ و واریانس σ^2 باشد ($X \sim N(\mu, \sigma^2)$) برای بدست آوردن احتمال $p(a < x < b)$ داریم:

$$\begin{aligned} p(a < x < b) &= p\left(\frac{a-\mu}{\sigma} < \frac{X-\mu}{\sigma}\right) \\ &= P\left(\frac{a-\mu}{\sigma} < Z < \frac{b-\mu}{\sigma}\right) \\ &= \Phi\left(\frac{b-\mu}{\sigma}\right) - \Phi\left(\frac{a-\mu}{\sigma}\right) \end{aligned}$$

۱۸-۷ مثال ۵: اگر $X \sim N(5, 16)$ X مطلوبست محاسبه $p(2 \leq X < 7)$.

حل: داریم $\mu = 5$, $\sigma^2 = 16$ پس:

$$\begin{aligned} p(2 \leq X < 7) &= p\left(\frac{2-5}{4} \leq \frac{X-5}{4} < \frac{7-5}{4}\right) \\ &= p\left(-\frac{3}{4} \leq Z < \frac{1}{2}\right) \\ &= \Phi\left(\frac{1}{2}\right) - \Phi\left(-\frac{3}{4}\right) = 0.6915 = 0.2266 = 0.4649 \end{aligned}$$

۱۹-۷ مثال ۶: نمرات دانشجویان یک کلاس از توزیع نرمال با میانگین ۱۵ و واریانس ۴ پیروی می‌کند اگر بدانیم تمامی نمرات از ۱۰ بیشتر هستند احتمال اینکه نمرات بین ۱۲ تا ۱۶ باشد چقدر است؟

حل: متغیر تصادفی X نرمال می‌باشد ($X \sim N(15, 4)$) می‌بایستی احتمال شرطی زیر را محاسبه کنیم:

$$\begin{aligned}
 p(12 < X < 16 \mid X > 10) &= \frac{p(12 < X < 16, X > 10)}{p(X > 10)} \\
 &= \frac{p(12 < X < 16)}{p(X > 10)} = \frac{p(\frac{12-15}{2} < \frac{X-15}{2} < \frac{16-15}{2})}{p(\frac{X-15}{2} > \frac{10-15}{2})} \\
 &= \frac{p(-\frac{3}{2} < Z < \frac{1}{2})}{p(Z > -\frac{5}{2})} = \frac{\varphi(\frac{1}{2}) - \varphi(-\frac{3}{2})}{\varphi(\frac{5}{2})} = \frac{0.6915 - 0.0668}{0.9938} = 0.6285
 \end{aligned}$$

۲۰-۷.۶.۲ تقریب توزیع دو جمله‌ای به توزیع نرمال

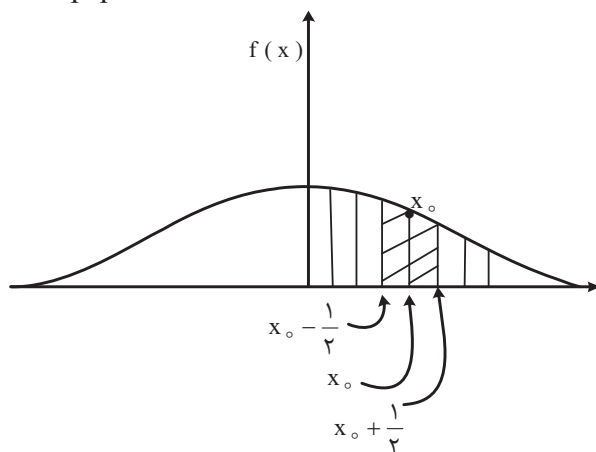
اگر در یک توزیع دو جمله‌ای تعداد آزمایشها بسیار زیاد باشد یعنی $n \rightarrow +\infty$ و در عین حال احتمال پیروزی و شکست تقریباً برابر باشند یعنی

$$p \approx q \approx \frac{1}{2}$$

در این صورت می‌توان مقادیر احتمال دو جمله‌ای را با استفاده از توزیع نرمال بدست آورد.

برای تقریب توزیع دو جمله‌ای به نرمال مقادیر امید و واریانس را برابر قرار می‌دهیم یعنی:

$$\mu = np, \quad \sigma^2 = npq$$



حال به شکل زیر توجه کنید:

$$p(X = x_0)$$

برای بدست آوردن

در توزیع دو جمله‌ای کافیست.

در توزیع نرمال مساحت مشخص شده در شکل را بدست بیاوریم زیرا می‌دانیم $p(X = x_0) = f(x_0)$. و از طرفی مساحت حاشور خورده تقریباً با مساحت یک ذوزنقه برابر است:

$$\begin{aligned}
 \text{مساحت ذوزنقه حاشور خورده} &= \text{ارتفاع} \times \frac{\text{مجموع دو قاعده}}{2} = (x_0 + \frac{1}{2} - (x_0 - \frac{1}{2})) \times \frac{f(x_0 + \frac{1}{2}) + f(x_0 - \frac{1}{2})}{2} \\
 &= \frac{f(x_0 + \frac{1}{2}) + f(x_0 - \frac{1}{2})}{2} \approx f(x_0)
 \end{aligned}$$

حال داریم:

$$p(X \leq x_0) = p(x_0 - \frac{1}{2} < X < x_0 + \frac{1}{2})$$

$$= P \frac{x_0 - \frac{1}{2} - \mu}{\sigma} < \frac{X - \mu}{\sigma} < \frac{x_0 + \frac{1}{2} - \mu}{\sigma}$$

$$= \varphi\left(\frac{x_0 + \frac{1}{2} - \mu}{\sigma}\right) - \varphi\left(\frac{x_0 - \frac{1}{2} - \mu}{\sigma}\right)$$

و با جاگذاری مقادیر $\mu = np$, $\sigma^2 = npq$ بدست می آوریم:

$$p(X=x_o) = \varphi\left(\frac{x_o + \frac{1}{2} - np}{\sqrt{npq}}\right) - \varphi\left(\frac{x_o - \frac{1}{2} - np}{\sqrt{npq}}\right)$$

به همین ترتیب:

$$p(X \leq x_o) = \varphi\left(\frac{x_o + \frac{1}{2} - np}{\sqrt{npq}}\right)$$

$$p(a < X \leq b) = \varphi\left(\frac{b + \frac{1}{2} - np}{\sqrt{npq}}\right) - \varphi\left(\frac{a - \frac{1}{2} - np}{\sqrt{npq}}\right)$$

۲۱-۷ مثال ۷: احتمال بارندگی در طول یک ماه در یک شهر $0/3$. مطلوبست احتمال اینکه حداقل در ۵۰ شهر از ۲۰۰ شهر کشور در طول یک

ماه آینده باران ببارد؟

حل: در این مثال X یک متغیر تصادفی دو جمله‌ای با پارامترهای $X \sim B(200, 0/3)$ می‌باشد می‌خواهیم احتمال $p(X > 50)$ را محاسبه کنیم:

$$p(X > 50) = 1 - p(X \leq 50) = 1 - \varphi\left(\frac{50 + \frac{1}{2} - (200 \times 0/3)}{\sqrt{200 \times 0/3 \times 0/7}}\right)$$

$$= 1 - \varphi\left(\frac{-9/5}{6/4}\right) = 1 - \varphi(-1/48) = 1 - 0/0694 = 0/9306$$

۲۲-۷ .۴ .۶ تقریب توزیع پواسون به توزیع نرمال

برای توزیعهای پواسون با λ بزرگ (بزرگتر یا مساوی ۵) می‌توان مقادیر احتمال توزیع پواسون را با استفاده از توزیع نرمال بدست آورد در این حالت روابط دقیقاً مشابه روابط تقریب دو جمله‌ای به نرمال می‌باشند با این تفاوت که به جای μ , σ^2 داریم:

$$p(X=x_o) = \varphi\left(\frac{x_o + \frac{1}{2} - \lambda}{\lambda}\right) - \varphi\left(\frac{x_o - \frac{1}{2} - \lambda}{\lambda}\right)$$

$$p(X \leq x_o) = \varphi\left(\frac{x_o + \frac{1}{2} - \lambda}{\lambda}\right)$$

$$p(a < X \leq b) = \varphi\left(\frac{b + \frac{1}{2} - \lambda}{\lambda}\right) - \varphi\left(\frac{a - \frac{1}{2} - \lambda}{\lambda}\right)$$

۲۳-۷ مثال ۸: تلفن یک شرکت در هر ساعت تقریباً ۲۰ مرتبه زنگ می‌زند مطلوبست احتمال اینکه تلفن در طول ۲ ساعت آینده حداقل ۲۰ و

حداکثر ۴۰ بار زنگ بزند؟

حل: X یک متغیر تصادفی پواسون می‌باشد واحد زمانی را هر یک ساعت در نظر می‌گیریم در این صورت $\lambda = 10$ می‌باشد و می‌بایستی احتمال $p(20 < X < 40)$ را محاسبه کنیم:

$$p(20 < X < 40) = \varphi\left(\frac{40 + \frac{1}{2} - 10}{10}\right) - \varphi\left(\frac{20 - \frac{1}{2} - 10}{10}\right)$$

$$= \varphi(3/05) - \varphi(0/95) = 0/9989 - 0/8289 = 0/17$$

فصل هشتم

توزیع های نمونه ای

در فصل های پیش نشان دادیم که با داشتن تابع توزیع یک متغیر تصادفی چگونه می توان مقدار احتمالات مورد نظر را بدست آورد. اما در عمل در بسیاری از اوقات می دانیم که یک متغیر تصادفی (مثلاً بارش باران در طول سال) از چه توزیع پیروی می کند اما مشکل اینجاست که مقدار پارامترهای توزیع را نمی دانیم. به عنوان مثال تعداد دفعات برنده شدن حساب بانکی یک شخص در طول ۱۰ سال یک متغیر تصادفی دو جمله ای با پارامترهای $X \sim B(10, p)$ می باشد. اما p یا همان احتمال برنده شدن حساب را نمی دانیم بنابراین برای بدست آوردن آن نیازمند روشهایی هستیم که بتوان پارامترهای مجهول را برآورد نمود در این فصل و فصل های بعدی روشهایی را برای این منظور ارائه می کنیم.

۸.۱ نمونه تصادفی

یک جمعیت آماری را با متغیر تصادفی X و توزیع احتمال $f_X(x)$ در نظر بگیرید. از این جمعیت n عضو را انتخاب می کنیم، این n عضو را یک نمونه n تایی گوئیم. اولین عضو از نمونه را با X_1 و دومین عضو را با X_2 و به همین ترتیب n امین عضو انتخابی را با X_n نمایش می دهیم. به این ترتیب نمونه n تایی به صورت $X_1, X_2, \dots, X_n$ خواهد بود که در آن هر یک از X_i یا خود یک متغیر تصادفی یا همان توزیع $f_X(x)$ می باشند. و از آنجا که تعداد اعضای جمعیت آماری نامتناهی یا بسیار زیاد می باشد ($N \rightarrow \infty$). در این صورت انتخاب X_i یا مستقل از X_j ($i \neq j$) می باشد. و این یعنی توزیع توام X_1 تا X_n برابر با حاصل ضرب توزیع X_1 تا X_n می باشد. حال با توجه به مطالب فوق تعریف نمونه تصادفی را بصورت زیر ارائه می کنیم:

تعریف:

گوئیم $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از متغیر تصادفی X می باشد اگر و فقط اگر تابع احتمال $f(X_1, X_2, \dots, X_n)$ بصورت زیر باشد:

$$f(X_1, X_2, \dots, X_n) = \prod_{i=1}^n f_X(x_i) \quad \forall (x_1, x_2, \dots, x_n)$$

حال در مثال های بعدی نحوه بکارگیری نمونه گیری را مشخص می کنیم.

مثال ۱: مدیر یک کارخانه کنسروسازی ادعا می کند که از هر ۱۰۰۰۰ کنسرو تولید شده تنها ۵ قوطی دارای اشکال می باشد، به وضوح برای اثبات ادعای وی نمی توانیم تمام قوطی ها را بررسی کنیم بنابراین از میان تمام قوطی ها ۱۰ قوطی انتخاب می کنیم (در این مثال برای سادگی تعداد کمتری نمونه انتخاب کردیم) و آنها را مورد بررسی قرار می دهیم. مطلوبست:

الف) تعریف متغیر تصادفی X و نمونه های تصادفی x_i $i=1, 2, \dots, 10$
 ب) تابع چگالی احتمال توام نمونه های تصادفی X_i .

حل: الف) متغیر تصادفی X را به این صورت تعریف می کنیم:

$$X = \begin{cases} 0 & \text{اگر قوطی سالم باشد} \\ 1 & \text{اگر قوطی ناسالم باشد} \end{cases}$$

به این ترتیب نمونه های تصادفی X_i به این صورت تعریف خواهد شد:

$$X = \begin{cases} 0 & \text{اگر نمونه } i \text{ ام سالم باشد} \\ 1 & \text{اگر نمونه } i \text{ ام ناسالم باشد} \end{cases}$$

بنابراین متغیر تصادفی X یک متغیر تصادفی برنولی با پارامتر p می باشد.

نمونه تصادفی عبارتست از $X_1, X_2, \dots, X_n$ که هر یک نیز به نوبه خود دارای توزیع برنولی با همان پارامتر p می باشند. یعنی داریم:

$$f_{X_i}(x_i) = p^{x_i} (1-p)^{1-x_i} \quad i=1, 2, \dots, 10$$

$$x_i = 0 \text{ یا } 1$$

$$= 0$$

برای سایر مقادیر

ب) از آنجا که نمونه ها از بین تعداد زیادی ($10/000$) کنسرو انتخاب می شوند بنابراین انتخاب با جایگذاری و بدون جایگذاری تقریباً معادل یکدیگر می باشند و می توان گفت که X_i ها از یکدیگر مستقل می باشند. در این صورت تابع چگالی توام X_i ها عبارتست از:

$$f(x_1, x_2, \dots, x_{10}) = \prod_{i=1}^{10} f(x_i) = p^{\sum_{i=1}^{10} x_i} (1-p)^{10 - \sum_{i=1}^{10} x_i}$$

مثال ۲: می‌دانیم که طول قد دانشجویان یک دانشگاه یک متغیر تصادفی نرمال با میانگین μ و واریانس σ^2 می‌باشد. اگر به تصادف ۱۰ نفر از دانشجویان انتخاب کنیم و X_1 را طول قد نفر اول و X_2 را طول قد نفر دوم و به همین ترتیب X_{10} را طول قد نفر دهم قرار دهیم. در این صورت مطلوبست:

الف) تابع چگالی احتمال توام X_1 تا X_{10} .

ب) احتمال اینکه قد هر ۱۰ نفر از ۱/۵ متر بیشتر باشد؟

ج) احتمال اینکه هیچ کس قدی بلندتر از ۲ متر نداشته باشد؟

حل: الف) از آنجا که هر یک از X_i ها نمونه‌های تصادفی می‌باشند که از یکدیگر مستقل بوده و از توزیع نرمال با میانگین μ و واریانس σ^2 پیروی می‌کند بنابراین:

$$f(X_1, X_2, \dots, X_n) = \prod_{i=1}^{10} f(x_i) = \prod_{i=1}^{10} \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{(x_i - \mu)^2}{2\sigma^2}} = \frac{1}{\sigma^{10} (2\pi)^5} e^{-\frac{\sum_{i=1}^{10} (x_i - \mu)^2}{2\sigma^2}}$$

ب) احتمال اینکه قد هر ۱۰ نفر از ۱/۵ متر بیشتر باشد برابر است با:

$$p(X_1 > 1/5) = 1 - p(X_1 \leq 1/5) = 1 - p\left(\frac{X_1 - \mu}{\sigma} < \frac{1/5 - \mu}{\sigma}\right) \\ = 1 - N_Z\left(Z < \frac{1/5 - \mu}{\sigma}\right)$$

برای X_1 تا X_{10} داریم:

$$p(X_1, \dots, X_{10} > 1/5) = \left(1 - N_Z\left(\frac{1/5 - \mu}{\sigma}\right)\right)^{10}$$

ج) مطابق بند (ب) مقدار این احتمال هم برابر است با:

$$p(X_1, \dots, X_{10} < 2) = \left(N_Z\left(\frac{2 - \mu}{\sigma}\right)\right)^{10}$$

همانطور که ملاحظه می‌کنید تابع چگالی توام و نتایج احتمالات به مقادیر μ و σ^2 وابسته می‌باشد. در فصل برآورد نشان می‌دهیم که با برآورد این پارامترها می‌توان مقادیر احتمالات را محاسبه نمود.

۲.۸ آماره‌ها

اگر با استفاده از عضوهای یک نمونه تصادفی (مثل $X_1, X_2, \dots$) یک تابع بسازیم که به مقادیر مجهول پارمتر وابستگی نداشته باشد به آن تابع یک آماره می‌گوییم.

به این ترتیب اگر $X_1, X_2, \dots, X_n$ یک نمونه تصادفی باشند توابع زیر همگی آماره می‌باشند:

$$X_2 + X_4 + X_8, \quad \frac{X_n}{X_1}, \quad \frac{1}{n} \sum_{i=1}^n (X_i)^2$$

$$\prod_{i=1}^n X_i, \quad \sum_{i=1}^n (X_i - 2)^2$$

اما توابع زیر بشرطی آماره می‌باشند که مقادیر μ و σ^2 مجهول نباشند:

$$X_n - \mu, \quad \frac{1}{n-1} \sum_1^n (X_i - \mu)^2, \quad \frac{X_2 - X_1}{\sqrt{\sigma^2}}, \quad \frac{X_1 - \mu}{\sigma}$$

به عنوان مثال $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ میانگین نمونه تصادفی می‌باشد که یک آماره هم می‌باشد. اما این میانگین با پارامتر μ در جامعه تفاوت می‌کند زیرا میانگین نمونه از نمونه‌ای به نمونه دیگر تغییر می‌ند اما پارامتر جامعه μ همواره ثابت است.

در استنباط آماری با گرفتن نمونه‌هایی از جامعه و بدست آوردن آماره‌هایی مثل $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ یا $S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$ می‌خواهیم

نتیجه‌گیری‌هایی در مورد پارامترهای جامعه مثل میانگین μ و واریانس σ^2 داشته باشیم.

حال همانطور که برای یک متغیر تصادفی X مقادیری مثل میانگین، واریانس، k امین گشتاور را معرفی نمودیم برای متغیر تصادفی X نیز این مقادیر را می‌توانیم تعریف کنیم.

مقدار $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ میانگین نمونه گفته می‌شود. به همین ترتیب مقدار $S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$ نیز مقدار واریانس نمونه و $S = \sqrt{S^2}$

نیز مقدار انحراف از معیار نمونه نامیده می‌شود.

k امین گشتاور نمونه عبارتست از:

$$M_K = \frac{1}{n} \sum_{i=1}^n X_i^k \quad k=1, 2, \dots$$

توجه کنید که مقادیر $\bar{X}$ و S^2 یک آماره می‌باشد و در نتیجه خود یک متغیر تصادفی می‌باشد این مطلب برای مقادیر مختلف K در M_K (مثل $M_1, M_2, \dots$) نیز برقرار می‌باشد.

مثال ۳: نشان دهید روابط زیر دارای S^2 برقرار است:

$$S^2 = \frac{1}{n-1} (\sum X_i^2 - n \bar{X}^2) \quad \text{(الف)}$$

$$S^2 = \frac{n}{n-1} (M_2 - \bar{X}^2) \quad \text{(ب)}$$

حل: الف)

$$\begin{aligned} S^2 &= \frac{1}{n-1} \sum_1^n (X_i - \bar{X})^2 = \frac{1}{n-1} \sum_1^n (X_i^2 - 2X_i \bar{X} + \bar{X}^2) \\ &= \frac{1}{n-1} \left(\sum_1^n X_i^2 - 2\bar{X} \sum_1^n X_i + n\bar{X}^2 \right) \\ &= \frac{1}{n-1} \left(\sum_1^n X_i^2 - 2\bar{X} (n\bar{X}) + n\bar{X}^2 \right) = \frac{1}{n-1} \left(\sum_1^n X_i^2 - n\bar{X}^2 \right) \end{aligned}$$

ب) بوضوح از بند الف داریم:

$$\begin{aligned} S^2 &= \frac{1}{n-1} \left(\sum_1^n X_i^2 - n\bar{X}^2 \right) = \frac{1}{n-1} \times n \left(\frac{1}{n} \sum_1^n X_i^2 - \bar{X}^2 \right) \\ &= \frac{n}{n-1} (M_2 - \bar{X}^2) \end{aligned}$$

در بخش توزیع‌های نمونه‌ای علت بکار بردن $\frac{1}{n-1}$ به جای $\frac{1}{n}$ را در محاسبه S^2 توضیح می‌دهیم.

مثال ۴: در مثال ۲ نتیجه، حاصل از نمونه‌گیری از قد ۱۰ نفر از دانشجویان بصورت زیر می‌باشد:

۱۷۸ , ۱۸۷ , ۱۶۱ , ۱۵۰ , ۱۷۰ , ۱۷۵ , ۱۹۲ , ۱۶۶ , ۱۷۳ , ۱۸۴

مطلوبست:

الف) میانگین و واریانس نمونه.

ب) گشتاور اول و دوم نمونه.

حل: الف)

$$\bar{X} = \frac{1}{10} \sum_{i=1}^{10} X_i = \frac{1}{10} (178 + 187 + 161 + 150 + 170 + 175 + 192 + 166 + 173 + 184) = 173/6$$

برای محاسبه واریانس نمونه از گشتاور اول و دوم استفاده می‌نیم:

$$M_1 = M_1 = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X} = 173/6$$

$$M_2 = \frac{1}{n} \sum_{i=1}^n X_i^2 = \frac{1}{10} \sum_{i=1}^{10} (X_i)^2 = \frac{1}{10} (178^2 + 187^2 + 161^2 + 150^2 + 170^2 + 175^2 + 192^2 + 166^2 + 173^2 + 184^2) = 30280/4$$

بنابراین S^2 برابر است با:

$$S^2 = \frac{n}{n-1} (M_2 - \bar{X}^2) = \frac{10}{9} (30280/4 - 173^2/6^2)$$

$$= \frac{10}{9} (143/44) = 159/37$$

$$\Rightarrow S = \sqrt{159/37} = 12/62$$

انحراف از معیار نمونه:

۸. ۳ آماره مرتب، میانه و برد نمونه

نمونه تصادفی $X_1, X_2, \dots, X_n$ از متغیر تصادفی X را در نظر بگیرید اگر مقادیر مشاهده شده نمونه تصادفی را از کوچک به بزرگ مرتب کنیم بطوریکه $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$ را داشته باشیم در این صورت $X_{(i)}$ را آماره مرتب می‌نامیم. $(i = 1, 2, \dots, n)$ همچنین به $X_{(1)}$ حداقل مقدار نمونه و به $X_{(n)}$ حداکثر مقدار نمونه گفته می‌شود. میانه نمونه عبارتست از مقدار وسط آماره مرتب به عبارت دقیقتر m_o را میانه نمونه در نظر می‌گیریم که بصورت زیر بدست می‌آید:

$$\text{اگر } n \text{ فرد باشد } m_o = X_{(\frac{n+1}{2})}$$

$$\text{اگر } n \text{ زوج باشد } = \frac{X_{(\frac{n}{2})} + X_{(\frac{n}{2}+1)}}{2}$$

برد نمونه نیز عبارتست از تفاضل بزرگترین آماره مرتب از کوچکترین آماره مرتب که با R_o نمایش داده می‌شود.

$$R_o = X_{(n)} - X_{(1)}$$

مثال ۵: نمونه تصادفی مشاهده شده در مثال ۴ را در نظر بگیرید، آماره مرتب، حداقل مقدار نمونه، حداکثر مقدار نمونه، میانه و برد نمونه را بدست

بیاورید؟

حل: با مرتب نمودن نمونه‌های مشاهده شده داریم:

$$X_{(1)}=150, \quad X_{(2)}=161, \quad X_{(3)}=166, \quad X_{(4)}=170, \quad X_{(5)}=173 \\ X_{(6)}=175, \quad X_{(7)}=178, \quad X_{(8)}=184, \quad X_{(9)}=187, \quad X_{(10)}=192$$

$X_{(1)}=150$ حداقل مقدار نمونه.

$X_{(10)}=192$ حداکثر مقدار نمونه.

$$\text{میانۀ نمونه} : \frac{X_{(\frac{10}{2})} + X_{(\frac{10}{2}+1)}}{2} = \frac{X_{(5)} + X_{(6)}}{2} = \frac{173 + 175}{2} = 174$$

$$\text{برد نمونه} : X_{(10)} - X_{(1)} = 192 - 150 = 42$$

۸-۱۱

۸. ۴. تابع توزیع نمونه‌ای

برای نمونه تصادفی $X_1, X_2, \dots, X_n$ و آماره مرتب $X_{(1)} \leq X_{(2)} \leq \dots \leq X_{(n)}$ تابع توزیع نمونه‌ای بصورت زیر تعریف می‌شود:

$$F^*(t) = \begin{cases} 0 & t < X_{(1)} \text{ برای} \\ \frac{r}{n} & X_{(r)} \leq t < X_{(r+1)} \quad (r = 1, \dots, n) \\ 1 & t > X_{(n)} \end{cases}$$

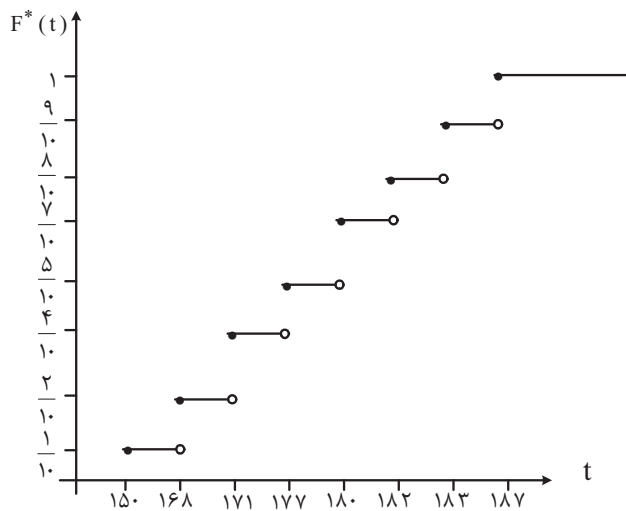
۸-۱۲

مثال ۶: اگر نمونه‌های مشاهده شد از قد دانشجویان در مثال ۴ بصورت زیر باشد مطلوبست تابع توزیع نمونه‌ای برای نمونه تصادفی مشاهده شده؟

$$150, 168, 171, 171, 177, 180, 180, 182, 183, 187$$

حل:

$$F^*(t) = \begin{cases} 0 & t < 150 \\ \frac{1}{10} & 150 \leq t < 168 \\ \frac{2}{10} & 168 \leq t < 171 \\ \frac{4}{10} & 171 \leq t < 177 \\ \frac{5}{10} & 177 \leq t < 180 \\ \frac{7}{10} & 180 \leq t < 182 \\ \frac{8}{10} & 182 \leq t < 183 \\ \frac{9}{10} & 183 \leq t < 187 \\ 1 & 187 \leq t \end{cases}$$



نمودار تابع توزیع نمونه‌ای فوق نیز بصورت زیر می‌باشد.

۸-۱۳

۵. نامساوی چبی چف

فرض کنید X یک متغیر تصادفی باشد که نحوه توزیع و تابع چگالی آنرا نمی‌دانیم در این صورت برای اظهار نظر در مورد مقادیر احتمال که متغیر تصادفی X قبول می‌کند می‌توانیم از قضیه چبی چف استفاده کنیم. این قضیه بصورت زیر می‌باشد:

اگر X یک متغیر تصادفی باشد و $g(X)$ یک تابع حقیقی غیر منفی از X باشد در این صورت:

$$p[g(X) \geq K] \leq \frac{1}{K} E[g(X)] \quad g(X) \geq 0, K > 0$$

با انتخاب $g(X) = (X - \mu_X)^2$ و $k^2 \sigma_X^2$ بجای K داریم:

$$p((X - \mu_X)^2 \geq k^2 \sigma_X^2) \leq \frac{1}{k^2 \sigma_X^2} E[(X - \mu_X)^2]$$

$$p(|X - \mu_X| \geq K \sigma_X) \leq \frac{1}{k^2}$$

بنابراین:

$$p(|X - \mu_X| < K \sigma_X) \geq 1 - \frac{1}{k^2}$$

مقادیر احتمال که توسط قانون چبی چف بدست می‌آیند بسته به مقدار k از دقت‌های مختلفی برخوردارند بطوریکه هر چه k بیشتر باشد دقت محاسبه احتمال توسط قانون چبی چف بیشتر است. این مطلب را در مثال زیر نشان می‌دهیم:

مثال ۷: اگر X یک متغیر تصادفی نرمال باشد در این صورت مطلوبست محاسبه احتمال $p(|X - \mu_X| < k \sigma_X)$ با استفاده از توزیع نرمال و قانون چبی چف به ازای مقادیر $k = 1, 2, 3, 4$.

حل: ابتدا مقدار احتمال را برای کتغیر تصادفی نرمال X بدست می‌آوریم:

$$p(|X - \mu_X| < K \sigma_X) = p\left(-k < \frac{X - \mu_X}{\sigma_X} < k\right)$$

$$= N_Z(k) - N_Z(-k) = 2N_Z(k) - 1$$

حال در جدول زیر مقادیر مختلف k را با استفاده از رابطه فوق و قانون چبی چف محاسبه می‌کنیم:

K	$2N_Z(k) - 1$	$1 - \frac{1}{k^2}$
۱	۰/۶۸۲۶	۰
۲	۰/۹۵۴۶	۰/۷۵
۳	۰/۹۹۷۴	۰/۸۸۸۹
۴	۱/۰۰۰	۰/۹۳۳۴

همانطور که ملاحظه می‌کنید مقادیر بدست آمده برای $k=1, 2$ از خطای بالایی برخوردار هستند اما با افزایش k مقدار بدست آمده از قانون چبی چف به مقدار واقعی نزدیک‌تر می‌شود. می‌توان تابع چگالی احتمال را طوری تعریف نمود که به ازای هر مقدار $k \geq 1$ مقدار واقعی احتمال با احتمال بدست آمده از رابطه چبی چف برابر باشد. این حالت را در مثال بعد نشان می‌دهیم:

مثال ۸: فرض کنید تابع چگالی احتمال متغیر تصادفی گسسته X بصورت زیر باشد:

$$f_X(x) = \frac{k^2 - 1}{k^2} \quad \text{برای } x = 0 \quad k \geq 1$$

$$= \frac{1}{2k^2} \quad \text{برای } x = -k, k$$

$$= 0 \quad \text{سایر مقادیر}$$

مطلوبست: الف) محاسبه μ_X , σ_X .

ب) محاسبه احتمال $p(|X - \mu_X| < k \sigma_X)$ به وسیله تابع احتمال و بوسیله قانون چبی چف و مقایسه آندو.

حل: الف)

$$\mu_X = \sum_{i=1}^n f(x_i) x_i = 0 \times \frac{k^2 - 1}{k^2} + k \left(\frac{1}{2k^2} \right) + (-k) \left(\frac{1}{2k^2} \right) = 0 + \frac{1}{2k} - \frac{1}{2k} = 0$$

$$E[X^2] = \sum_{i=1}^n x_i^2 f(x_i) = 0 \times \frac{k^2 - 1}{k^2} + k^2 \left(\frac{1}{2k^2} \right) + (-k)^2 \left(\frac{1}{2k^2} \right)$$

$$= 0 + \frac{1}{2} + \frac{1}{2} = 1 \quad \sigma_X^2 = E[X^2] - E^2[X] = \sigma_X^2 = 1 - 0^2 = 1$$

ب) ابتدا احتمال دقیق را محاسبه می‌کنیم:

$$\mu_X = 0 \quad \sigma_X = 1$$

$$\Rightarrow p(|X - \mu_X| < k \sigma_X) = p(|X| < K)$$

$$p(-k < X < k) = p(X = -k) + p(X = k) = \frac{1}{2k^2} + \frac{1}{2k^2} = \frac{1}{k^2}$$

اما طبق قانون چبی چف رابطه زیر برقرار می‌باشد:

$$p(|X - \mu_X| < k \sigma_X) \leq \frac{1}{k^2}$$

مقدار $\frac{1}{k^2}$ دقیقاً برابر با احتمال محاسبه شده می‌باشد بنابراین در این حالت مخصوص مقدار دقیق احتمال با مقدار بدست آمده از قانون چبی چف برابر می‌باشد و به عبارتی حالت مساوی در قانون چبی چف رخ می‌دهد.

همانطور که قبلاً اشاره کردیم تابع نرمال نقش مهمی در احتمالات بازی می‌کند زیرا اولاً محاسبه احتمال نرمال با استفاده از تابع نرمال استاندارد و جدول مربوط به آن بسیار ساده می‌باشد ثانیاً بسیاری از محاسبات احتمال را همانطور که در فصلهای قبل نشان دادیم تحت شرایط خاصی می‌توان با استفاده از تابع نرمال تقریب زد. حال در اینجا نشان می‌دهیم که اگر متغیرهای تصادفی X_i ($i = 1, 2, \dots$) از طکدیگر مستقل باشند و از یک تابع توزیع پیروی کنند، مجموع آنها به سمت تابع نرمال میل می‌کند. در اینجا نمی‌دانیم که متغیرهای تصادفی X_i از چه توزیعی پیروی می‌کند. صورت قضیه حد مرکزی بصورت زیر می‌باشد:

فرض کنید $X_1, X_2, \dots, X_n$ نمونه‌ای تصادفی از جامعه نامتناهی X با میانگین μ و σ^2 باشند در این صورت دنباله متغیرهای تصادفی $Z_1, Z_2, \dots, Z_n$ را به صورت زیر تعریف می‌کنیم:

$$Z_n = \frac{\bar{X}_n - \mu}{\frac{\sigma}{\sqrt{n}}} \quad n = 1, 2, \dots$$

که در آن $\bar{X}_n = \frac{1}{n} \sum_{i=1}^n X_i$ در این صورت اگر $F_{Z_n}(t)$ تابع توزیع متغیرهای تصادفی Z_n به ازای مقدار حقیقی t باشد داریم:

$$\lim_{n \rightarrow \infty} F_{Z_n}(t) = N_Z(t)$$

که در آن $N_Z(t)$ تابع توزیع نرمال استاندارد می‌باشد.

به عبارت ساده‌تر متغیر تصادفی $Z = \frac{\bar{X}_n - \mu}{\frac{\sigma}{\sqrt{n}}}$ وقتی n به سمت بی‌نهایت میل می‌کند دارای توزیع نرمال استاندارد $N(0, 1)$ می‌باشد. یعنی

برای n های به اندازه کافی بزرگ داریم $F_{Z_n}(t) = N_Z(t)$ اما برای هر مقدار n داریم:

$$F_{\bar{X}_n}(t) = F_{Z_n}\left(\frac{t - \mu}{\frac{\sigma}{\sqrt{n}}}\right)$$

$$F_{\bar{X}_n}(t) = N_Z\left(\frac{t - \mu}{\frac{\sigma}{\sqrt{n}}}\right)$$

در نتیجه برای n های بزرگ:

رابطه فوق به این معنی است که تابع توزیع میانگین تعداد زیادی از متغیرهای تصادفی مستقل با توزیع یکسان تقریباً مساوی است با توزیع نرمال استاندارد. به همین دلیل است که تعداد زیادی از قوانین احتمال را بوسیله نرمال تقریب می‌زنند.

۸-۱۷

به عنوان مثال به تقریب دو جمله‌ای به نرمال توجه کنید که در فصل قبل روابط آنرا بصورت زیر نشان دادیم:

$$p(X = x_0) = N\left(\frac{x_0 + \frac{1}{2} - np}{\sqrt{npq}}\right) - N\left(\frac{x_0 - \frac{1}{2} - np}{\sqrt{npq}}\right)$$

می‌دانیم اگر متغیرهای تصادفی برنولی و مستقل $X_1, X_2, \dots, X_n$ را هر یک با پارامتر p در نظر بگیریم در این صورت $Y = \sum_{i=1}^n X_i$ یک متغیر

تصادفی دو جمله‌ای با پارامتر n و p می‌باشد حال طبق قضیه حد مرکزی میانگین X_i ها به توزیع نرمال میل می‌کند یعنی تابع توزیع

$$\frac{Y}{n} = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X}_n$$

وقتی n بزرگ باشد برابر است با

$$F_{\frac{Y}{n}}(t) = N_Z\left(\frac{t-\mu}{\frac{\sigma}{\sqrt{n}}}\right)$$

که در آن امید ریاضی $\frac{Y}{n}$ برابر است با $\mu = p$ و واریانس $\frac{Y}{n}$ نیز برابر با $\sigma^2 = pq$ می باشد حال داریم:

$$F_{\frac{Y}{n}}(t) = N_Z\left(\frac{t-\mu}{\frac{\sigma}{\sqrt{n}}}\right) = N_Z\left(\frac{t-p}{\sqrt{\frac{pq}{n}}}\right)$$

حال $F_{\frac{Y}{n}}(t)$ را به $F_Y(t)$ تبدیل می کنیم:

$$F_Y(t) = F_{\frac{Y}{n}}\left(\frac{t}{n}\right) = N_Z\left(\frac{\frac{t}{n}-p}{\sqrt{\frac{pq}{n}}}\right) = N_Z\left(\frac{t-np}{\sqrt{npq}}\right)$$

در نتیجه $F_Y(t) = N_Z\left(\frac{t-np}{\sqrt{npq}}\right)$ که با استفاده از روش ذوزنقه ای که در فصل قبل نشان دادیم می توان مقدار احتمال $p(X = x_o)$ را بصورت زیر بدست آورد:

$$p(X = x_o) = N\left(\frac{x_o + \frac{1}{2} - np}{\sqrt{npq}}\right) - N\left(\frac{x_o - \frac{1}{2} - np}{\sqrt{npq}}\right)$$

فصل نهم

توزیع های برخی از آماره ها

در فصل قبل با آماره ها و تعاریف و مقدمات آنها آشنا شدید، نشان دادیم که آماره ها خود متغیر می باشند و می توان تابع توزیع آنها را معنی نمود، حال در این فصل چندین آماره خاص را مورد بررسی قرار می دهیم و میانگین و واریانس آنها را محاسبه می کنیم.

۱.۹ توزیع نمونه ای میانگین نمونه $\bar{X}$

از یک جامعه با میانگین μ و واریانس δ^2 تعداد n نمونه $X_1, X_2, \dots, X_n$ انتخاب می کنیم، آماره $\bar{X} = \frac{1}{n} \sum_{i=1}^n X_i$ را در نظر بگیرید

می خواهیم میانگین و واریانس $\bar{X}$ را به عنوان یک متغیر تصادفی بدست بیاوریم. طبق قضیه حد مرکزی اگر مقدار n به اندازه کافی بزرگ اختیر شود متغیر تصادفی $Z = \frac{\bar{X} - \mu}{\frac{\delta}{\sqrt{n}}}$ یک متغیر تصادفی نرمال استاندارد می باشد. از طرفی داریم:

$$\bar{X} = \frac{\delta}{\sqrt{n}} Z + \mu, \quad E[Z] = 0, \quad \delta_Z^2 = 1$$

$$E[\bar{X}] = \frac{\delta}{\sqrt{n}} E[Z] + \mu = \frac{\delta}{\sqrt{n}} \times 0 + \mu = \mu$$

بنابراین:

بنابراین آماره $\bar{X}$ به عنوان یک متغیر تصادفی اگر تعداد نمونه ها زیاد باشد ($n > 30$) همواره دارای میانگین μ یا همان میانگین جامعه می باشد.

۹-۳

مثال ۱: در نقطه ای از یک رودخانه در هر دقیقه بطور متوسط تعداد ۲۰ ماهی عبور می کند بدیهی است که ماهیگیر تنها در صورتی می تواند در هر دقیقه این تعداد ماهی را صید کند که تقریباً تمامی طول رودخانه را با تور ماهیگیری پوشش دهد یا به عبارتی در هر دقیقه از بیشتر نقاط در طول رودخانه نمونه برداری کند که معادل است با انتخاب تعداد زیادی نمونه. حال واریانس $\bar{X}$ را محاسبه می کنیم:

$$\text{var}(\bar{X}) = \text{var}\left(\frac{\delta}{\sqrt{n}} Z + \mu\right) = \text{var}\left(\frac{\delta}{\sqrt{n}} Z\right) = \frac{\delta^2}{\sqrt{n}} \text{var}(Z) = \frac{\delta^2}{\sqrt{n}}$$

بنابراین واریانس $\bar{X}$ برابر با $\frac{\delta^2}{\sqrt{n}}$ می باشد.

۹-۴

مثال ۲: محصول شیر تولیدی یک کارخانه بطور متوسط ۲۰ روز پس از تاریخ تولید قابل مصرف می باشد که با واریانس ۴ روز رخ می دهد. می خواهیم از این کارخانه ۱۰ پاکت شیر تهیه کنیم. احتمال اینکه این پاکتهای شیر بطور متوسط حداقل ۱۸ روز قابل مصرف باشند چقدر است؟

حل: برای به دست آوردن احتمال می بایستی $p(\bar{X} \geq 18)$ را محاسبه کنیم که در واقع از تابع توزیع متغیر تصادفی $\bar{X}$ که در اینجا نرمال μ $\frac{\delta^2}{\sqrt{n}}$

در نظر گرفته می شود استفاده می کنیم. بنابراین:

$$\mu_{\bar{X}} = 20, \quad n = 10$$

$$\delta_{\bar{X}}^2 = \frac{\delta^2}{n} = \frac{(4)^2}{10} = \frac{16}{10} \rightarrow \frac{\delta}{\bar{X}} = \sqrt{\frac{16}{10}} = 1/26$$

$$\begin{aligned}
p(\bar{X} \geq 18) &= p\left(\frac{\bar{X} - \mu_{\bar{X}}}{\delta_{\bar{X}}} \geq \frac{18 - \mu_{\bar{X}}}{\delta_{\bar{X}}}\right) \\
&= p\left(\frac{\bar{X} - \mu_{\bar{X}}}{\delta_{\bar{X}}} \geq \frac{18 - 20}{1/\sqrt{6}}\right) = p\left(\frac{\bar{X} - \mu_{\bar{X}}}{\delta_{\bar{X}}} \geq 1/\sqrt{6}\right) \\
p &= (Z \geq -1/\sqrt{6}) = 1 - p(Z \leq -1/\sqrt{6}) \\
&= 1 - N_Z(-1/\sqrt{6}) = 1 - 0.571 = 0.429
\end{aligned}$$

۹-۵-۱

۹-۲ توزیع نمونه‌ای واریانس نمونه S^2 .

حال به محاسبه توزیع نمونه‌ای واریانس S^2 و بدست آوردن میانگین و واریانس توزیع آن می‌پردازیم.

همانطور که از فصل‌های قبل به یاد دارید واریانس نمونه‌ها با S^2 نشان داده می‌شود که بصورت زیر تعریف می‌شود:

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2$$

قبل محاسبه توزیع نمونه‌ای S^2 ابتدا میانگین و واریانس آنرا (از آنجا که محاسبه آن ساده می‌باشد) بدست می‌آوریم:

می‌دانیم واریانس $\bar{X}$ برابر $\frac{\delta^2}{n}$ می‌باشد بنابراین:

$$\begin{aligned}
\text{var}(\bar{X}) &= \frac{\delta^2}{n} \rightarrow E[(\bar{X} - \mu)^2] = \frac{\delta^2}{n} \\
E[\bar{X}^2] - E^2[X] &= E[\bar{X}^2] - \mu^2 = \frac{\delta^2}{n} \Rightarrow E[\bar{X}^2] = \frac{\delta^2}{n} + \mu^2
\end{aligned}$$

حال داریم:

$$\begin{aligned}
E[S^2] &= E\left[\frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2\right] = \frac{1}{n-1} E\left[\sum_{i=1}^n (X_i - \bar{X})^2\right] \\
&= \frac{1}{n-1} E\left[\sum_{i=1}^n X_i^2 - 2\bar{X} \sum_{i=1}^n X_i + \sum_{i=1}^n \bar{X}^2\right] \\
&= \frac{n}{n-1} E\left[\frac{1}{n} \sum_{i=1}^n X_i^2 - 2\bar{X} \frac{1}{n} \sum_{i=1}^n X_i + \frac{1}{n} \sum_{i=1}^n \bar{X}^2\right] \\
&= \frac{n}{n-1} E\left[\frac{1}{n} \sum_{i=1}^n X_i^2 - 2\bar{X}^2 + \bar{X}^2\right] \\
&= \frac{n}{n-1} E\left[\frac{1}{n} \sum_{i=1}^n X_i^2 - \bar{X}^2\right] = \frac{n}{n-1} \left(E\left[\frac{1}{n} \sum_{i=1}^n X_i^2\right] - E[\bar{X}^2]\right) \\
&= \frac{n}{n-1} \left(E[M_2] - \left(\frac{\delta^2}{n} + \mu^2\right)\right)
\end{aligned}$$

۹-۵-۲

که در آن M_2

گشتاور دوم نمونه‌ها می‌باشد و داریم: $E[M_2] = m_2$ که m_2 نیز گشتاور دوم متغیر تصادفی X می‌باشد. در نهایت داریم:

$$= \frac{1}{n-1} (m_2 - \mu^2 - \frac{\delta^2}{n})$$

$$m_2 = \delta^2 + \mu^2 \quad \text{اما:}$$

$$= \frac{n}{n-1} (\delta^2 + \mu^2 - \mu^2 - \frac{\delta^2}{n}) = \frac{n}{n-1} (\delta^2 - \frac{\delta^2}{n}) = \delta^2$$

به نتیجه جالبی رسیدیم و آن اینکه میانگین واریانس نمونه‌ها برابر با واریانس جامعه می‌باشد در واقع علت اینکه در محاسبه S^2 از $\frac{1}{n-1}$ استفاده

می‌کنیم این است که در محاسبات فوق در نهایت جواب ساده شده و δ^2 به عنوان میانگین S^2 ($E[S^2] = \delta^2$) بدست می‌آید. در حالی که اگر

$$S^2 \text{ را بصورت زیر تعریف می‌کردیم } S^2 = \frac{1}{n} \sum_{i=1}^n (X_i - \bar{X})^2 \text{ جواب نهایی } E[S^2] \text{ برابر با } \sigma^2 \left(\frac{n}{n-1} \right) \text{ می‌شد که بلیل وجود } n \text{ در جواب}$$

نتیجه مطلوبی نیست. حال به جای بدست آوردن توزیع S^2 توزیع $S^2 \left(\frac{n-1}{\delta^2} \right)$ را بدست می‌آوریم زیرا در محاسبه مقدار میانگین بصورت زیر

ساده‌تر می‌شود.

$$E \left[\frac{n-1}{\delta^2} - S^2 \right] = \frac{n-1}{\delta^2} E[S^2] = \frac{n-1}{\delta^2} \delta^2 = n-1$$

برای بدست آوردن توزیع $S^2 \left(\frac{n-1}{\delta^2} \right)$ ابتدا توزیع χ^2 -دو یا مربع کای را معرفی می‌کنیم:

۹-۶

۱.۲.۹ توزیع χ^2 دو

فرض کنید Z یک متغیر تصادفی نرمال استاندارد باشد $Z \sim N(0, 1)$ در این صورت متغیر تصادفی $W = Z^2$ را یک متغیر تصادفی χ^2 (خی)

(دو) با یک درجه آزادی می‌نامیم و به صورت χ^2_1 نمایش می‌دهیم که تابع توزیع آن بصورت زیر بدست می‌آید:

قبلاً نشان دادیم که اگر $Y = X^2$ باشد تابع توزیع $F_Y(t)$ برابر است با:

$$F_Y(t) = F_X(\sqrt{t}) - F_X(-\sqrt{t}) \quad t \geq 0$$

$$= 0 \quad t < 0$$

حال داریم: $W = Z^2$

$$F_W(t) = F_Z(\sqrt{t}) - F_Z(-\sqrt{t})$$

اما: $F_Z(\sqrt{t}) = N_Z(t)$ و $N_Z(-a) = 1 - N_Z(a)$ بنابراین:

$$F_W(t) = N_Z(\sqrt{t}) - [1 - N_Z(\sqrt{t})] = 2N_Z(\sqrt{t}) - 1$$

بنابراین با داشتن جدول مقادیر توزیع نرمال استاندارد به راحتی می‌توان مقادیر توزیع χ^2_1 را به دست آورد. مقدار دقیق تابع چگالی χ^2_1 بصورت زیر می‌باشد:

$$f_X(x) = \frac{1}{\sqrt{2\pi}} x^{-\frac{1}{2}} e^{-\frac{x}{2}} \quad x > 0$$

میانگین، واریانس و تابع مولد گشتاور آن عبارتند از:

$$E[W] = 1 \quad \text{var}(W) = 2 \quad m_W(t) = \left(\frac{1}{1-2t} \right)^{\frac{1}{2}}$$

۹-۷

حال فرض کنید $X_1, X_2, \dots, X_n$ نمونه‌هایی تصادفی از یک متغیر تصادفی نرمال X با میانگین μ و واریانس δ^2 باشند در این صورت:

$$Y = \sum_{i=1}^n \frac{(X_i - \mu)^2}{\delta^2}$$

یک متغیر تصادفی χ^2 با n درجه آزادی نامیده می‌شود و با χ_n^2 نمایش داده می‌شود. به بیانی دیگر اگر $Z_1, Z_2, \dots, Z_n$ تعداد n متغیر تصادفی نرمال استاندارد $N(0, 1)$ باشند در این صورت $Y = Z_1^2, Z_2^2, \dots, Z_n^2$ یک متغیر تصادفی χ^2 با n درجه آزادی است. تابع چگالی متغیر تصادفی χ_n^2 بصورت زیر می‌باشد:

$$f_Y(y) = \frac{1}{\Gamma(\frac{n}{2})} \left(\frac{1}{2} \right)^{\frac{n}{2}} y^{\frac{n}{2}-1} e^{-\frac{y}{2}} \quad y > 0$$

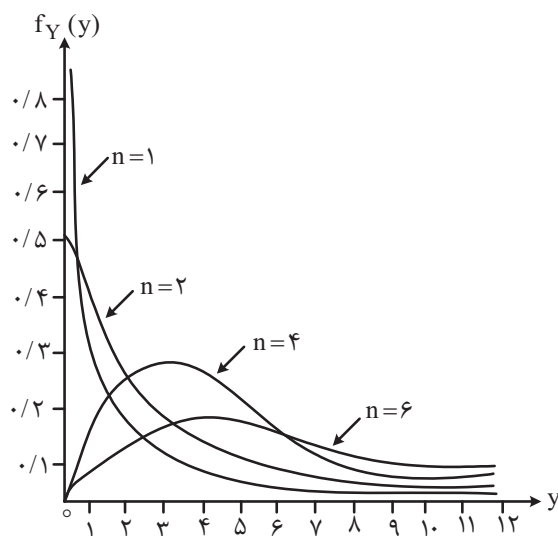
توجه کنید که متغیر تصادفی χ_n^2 حالت خاصی از متغیر تصادفی گاما می‌باشد که در آن $k = \frac{n}{2}$ و $\lambda = \frac{1}{2}$ می‌باشد.

مقادیر میانگین، واریانس و تابع مولد گشتاور متغیر تصادفی χ_n^2 بصورت زیر می‌باشد:

$$E[Y] = n, \quad \text{var}(Y) = 2n$$

$$m_Y(t) = \left(\frac{1}{1-2t} \right)^{\frac{n}{2}}$$

نمودار منحنی متغیر تصادفی χ_n^2 دارای نقطه ماکزیممی در $y = n-2$ می‌باشد و با افزایش n نمودار منحنی هر چه بیشتر به نمودار نزدیک نزدیکتر می‌شود این مطلب در شکل زیر نشان داده شده است:



نکته: اگر $X_1, X_2, \dots, X_n$ متغیرهای تصادفی خی دو به ترتیب با $r_1, r_2, \dots, r_n$ درجه آزادی باشند و همچنین دو بدو مستقل از یکدیگر باشند در این صورت متغیر تصادفی $Y = \sum_{i=1}^n X_i$ دارای توزیع خی دو با $R = \sum_{i=1}^n r_i$ درجه آزادی می‌باشد.

۹-۸

مثال ۳: فرض کنید X_1, X_2, X_3 سه متغیر تصادفی مستقل خی دو به ترتیب با ۷ و ۳ و ۱ درجه آزادی باشند، اگر $W = X_1 + X_2 + X_3$ مطلوبست میانگین و واریانس متغیرهای تصادفی W ؟

حل: متغیر تصادفی W یک متغیر تصادفی خی دو با $1 + 3 + 7 = 11$ درجه آزادی می‌باشد بنابراین میانگین و واریانس آن عبارتست از:

$$E[W] = 11$$

$$\text{var}(W) = 2 \times 11 = 22$$

حال به قضیه بسیار مهم توجه کنید:

قضیه: اگر $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از متغیر تصادفی نرمال X با میانگین μ و واریانس δ^2 باشند آنگاه:

الف) $\bar{X}$ و S^2 از یکدیگر مستقل هستند.

ب) $\frac{(n-1)}{\delta^2} S^2$ یک متغیر تصادفی χ^2 با $n-1$ درجه آزادی می‌باشد توجه کنید که $\frac{(n-1)}{\delta^2} S^2$ برابر است با:

$$\left(\frac{n-1}{\delta^2}\right) S^2 = \frac{n-1}{\delta^2} \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 = \sum_{i=1}^n \frac{(X_i - \bar{X})^2}{\delta^2}$$

و این با آنچه در تعریف متغیر تصادفی χ_n^2 متفاوت می‌باشد زیرا در تعریف از میانگین جامعه μ استفاده نمودیم: $\sum_{i=1}^n \frac{(X_i - \mu)^2}{\delta^2}$

۹-۹

حال نشان می‌دهیم که چرا با جایگزین نمودن $\bar{X}$ به جای μ یک درجه آزادی از n کم می‌شود و درجه آزادی $\frac{(n-1)}{\delta^2} S^2$ برابر با $n-1$ می‌شود.

$$\begin{aligned} \sum_{i=1}^n (X_i - \mu)^2 &= \sum_{i=1}^n [(X_i - \bar{X}) + (\bar{X} - \mu)]^2 \\ &= \sum_{i=1}^n (X_i - \bar{X})^2 + \sum_{i=1}^n (\bar{X} - \mu)^2 + 2(\bar{X} - \mu) \sum_{i=1}^n (X_i - \bar{X}) \\ &= \sum_{i=1}^n (X_i - \bar{X})^2 + n(\bar{X} - \mu)^2 \end{aligned}$$

در نتیجه بدست می‌آید:

$$\frac{\sum (X_i - \mu)^2}{\delta^2} = \frac{\sum (X_i - \bar{X})^2}{\delta^2} + \frac{(\bar{X} - \mu)^2}{\frac{\delta^2}{n}}$$

حال می‌دانیم که عبارت سمت چپ تساوی فوق یک متغیر تصادفی X با n درجه آزادی است و عبارت $\frac{(\bar{X}-\mu)^2}{\frac{\delta^2}{n}}$ نیز یک متغیر تصادفی X^2

با ۱ درجه آزادی است بنابراین جمله $\frac{\sum (X_i - \mu)^2}{\delta^2}$ می‌بایستی یک متغیر تصادفی χ^2 با $n-1$ درجه آزادی باشد.

مثال ۴: از یک جامعه نرمال با میانگین μ و واریانس δ^2 یک نمونه تصادفی Y انتخاب می‌کنیم. متغیر تصادفی $Y = \frac{10 S^2}{\delta^2}$ را در نظر

بگیرید مطلوبست:

الف) محاسبه $p(3/94 < Y < 18/3)$.

ب) با استفاده از احتمال فوق رابطه بین S^2 و δ^2 را در جامعه معین کنید.

حل: الف) ابتدا احتمال را بصورت زیر ساده می‌کنیم:

$$p(3/94 < Y < 18/3) = p(Y < 18/3) - p(Y < 3/94)$$

حال از روی جدول مقادیر احتمال برای متغیر تصادفی Y با دو و با توجه به اینکه $Y = \frac{10 S^2}{\delta^2}$ یک متغیر تصادفی χ^2 با ۱۰ درجه آزادی است

مقادیر احتمال بصورت زیر بدست می‌آیند:

$$p(Y < 18/3) - p(Y < 3/94) = 0.95 / 0.05 = 0.9$$

ب) می‌دانیم $p(3/94 < Y < 18/3) = 0.9$ بنابراین:

$$p(3/94 < Y = \frac{10 S^2}{\delta^2} < 18/3) = p\left(\frac{3/94 \delta^2}{10} < S^2 < \frac{18/3 \delta^2}{10}\right)$$

$$\Rightarrow p(0.394 \delta^2 < S^2 < 1.83 \delta^2) = 0.9$$

حال می‌بینیم که S^2 در بازه‌ای نزدیک δ^2 قرار دارد یعنی اگر از یک جامعه نرمال با میانگین μ و واریانس δ^2 تعداد ۱۱ نمونه انتخاب کنیم به احتمال بسیار زیاد (۰/۹) مقدار واریانس نمونه‌ها نزدیک به واریانس جامعه می‌باشد.

۹-۱۱-۱

۳.۹ تابع توزیع آماره مرتب

قبلاً آماره مرتب را معرفی نمودیم، می‌دانیم آماره‌های مرتب خود یک متغیر تصادفی می‌باشند. بنابراین می‌توانیم تابع توزیع آنها را بدست بیاوریم.

فرض کنید $X_1, X_2, \dots, X_n$ تعداد n نمونه تصادفی از متغیر تصادفی X باشند در این صورت تابع توزیع برای آماره مرتب $X_{(1)}, X_{(2)}, \dots, X_{(n)}$ عبارتست از:

$$F_{X(j)}(t) = \sum_{k=j}^n \binom{n}{k} (F_X(t))^k (1 - F_X(t))^{n-k} \quad t \in \mathbb{R}$$

$F_X(t)$ تابع توزیع متغیر تصادفی X می‌باشد.

$F_{X(j)}(t)$ به صورت زیر بدست می‌آید:

اگر حداقل j متغیر تصادفی نمونه، کمتر یا مساوی با t باشند در این صورت $X_{(j)}$ یا j امین آماره مرتب کمتر یا مساوی با t می‌باشد. تمام متغیرهای

تصادفی نمونه مستقل می‌باشند و برای تمام متغیرهای تصادفی نمونه $F_X(t)$ برابر است با:

احتمال کمتر یا مساوی با t بودن. بنابراین تعداد متغیرهای تصادفی نمونه که کمتر یا مساوی با t می‌باشند یک متغیر تصادفی دو جمله‌ای با پارامترهای n و $p = F_X(t)$ می‌باشد، در نتیجه:

$$f_{X(j)}(t) = \sum_{k=j}^n \binom{n}{k} (F_X(t))^{j-1} (1-F_X(t))^{n-j} f_X(t)$$

با قرار دادن $j = 1$ و $j = n$ تابع توزیع و چگالی برای $\min(X_j)$ و $\max(X_j)$ بصورت زیر بدست می‌آید:

$$\min(X_j) \text{ تابع توزیع } F_{X(1)}(t) = \sum_{k=1}^n \binom{n}{k} (F_X(t))^k (1-F_X(t))^{n-k} = 1 - (1-F_X(t))^n$$

$$\min(X_j) \text{ تابع چگالی } f_{X(1)}(t) = \left(\binom{n}{1} (f_X(t))^{1-1} (1-F_X(t))^{n-1} f_X(t) \right) = n (1-F_X(t))^{n-1} f_X(t)$$

همانطور که ملاحظه می‌کنید رابطه $f_X(t) = \frac{dF_{X(1)}(t)}{dt}$ نیز بین تابع توزیع و تابع چگالی برقرار است.

$$\max(X_j) \text{ تابع توزیع } F_{X(n)}(t) = \sum_{k=n}^n \binom{n}{k} (F_X(t))^k (1-F_X(t))^{n-k} = (F_X(t))^n$$

$$\max(X_j) \text{ تابع چگالی } f_{X(n)}(t) = n \binom{n}{n} (f_X(t))^{n-1} (1-F_X(t))^{n-n} f_X(t) = n (F_X(t))^{n-1} f_X(t)$$

مثال ۵: متغیر تصادفی X را مدت زمان بین هر دو تماس تلفنی در یک شرکت در نظر می‌گیریم که بطور میانگین هر ۱۰ دقیقه یکبار رخ می‌دهد

اگر بطور تصادفی ۴ نمونه از متغیر تصادفی X اختیار کنیم مطلوبست:

الف) تابع توزیع آماره مرتب نمونه تصادفی.

ب) تابع توزیع و چگالی ماکزیمم و مینیمم نمونه‌ها.

ج) احتمال اینکه در نمونه‌های مشاهده شده کوچکترین بازه زمانی بین دو تماس بیشتر از ۲ دقیقه باشد چقدر است؟

حل: الف) ابتدا توجه کنید که X یک متغیر تصادفی نمایی با پارامتر $\lambda = \frac{1}{10}$ می‌باشد. بنابراین تابع توزیع X عبارتست از:

$$F_X(t) = 1 - e^{-\lambda t} = 1 - e^{-\frac{1}{10}t} \quad t > 0$$

بنابراین تابع توزیع آماره مرتب برای ۴ نمونه مشاهده شده عبارتست از:

$$F_{X(j)}(t) = \sum_{k=j}^4 \binom{4}{k} (1 - e^{-\frac{1}{10}t})^k (e^{-\frac{1}{10}t})^{4-k}$$

۴-۱۲-۲

ب) ماکزیمم آماره‌های مرتب در این مثال $X_{(4)}$ می‌باشد. که تابع توزیع و چگالی آن بدست می‌آید:

$$F_{X(4)}(t) = (F_X(t))^4 = (1 - e^{-\frac{1}{10}t})^4 \quad t > 0$$

$$F_{X(4)}(t) = 4 (1 - e^{-\frac{1}{10}t})^3 \frac{1}{10} e^{-\frac{1}{10}t}$$

مینیمم آماره‌های مرتب $X_{(1)}$ می‌باشد که تابع توزیع و چگالی آن نیز بصورت زیر می‌باشد:

$$F_{X(1)}(t) = 1 - (1 - F_X(t))^4 = 1 - (1 - 1 + e^{-\frac{1}{10}t})^4 = 1 - e^{-\frac{4}{10}t}$$

$$F_{X(1)}(t) = 4(1 - 1 + e^{-\frac{1}{10}t})^3 \cdot \frac{1}{10} e^{-\frac{1}{10}t} = \frac{4}{10} e^{-\frac{3}{10}t} e^{-\frac{1}{10}t} = \frac{4}{10} e^{-\frac{4}{10}t}$$

توجه کنید که توزیع کوچکترین آماره مرتب در این مثال خود یک متغیر تصادفی نمایی با پارامتر $\lambda = \frac{4}{10}$ می باشد.

ج) احتمال اینکه کوچکترین بازه زمانی بین دو تماس در نمونه‌ها بیشتر از ۲ دقیقه باشد برابر است با: $p(X_{(1)} > 2)$ و داریم:

$$p(X_{(1)} > 2) = 1 - p(X_{(1)} < 2) = 1 - F_{X(1)}(t=2)$$

$$= 1 - (1 - e^{-\frac{4}{10} \times 2}) = e^{-0.8} \approx 0.449$$

۹-۱۳

مثال ۶: برای آماره مرتب $X_{(1)}, X_{(2)}, \dots, X_{(n)}$ از یک متغیر تصادفی پیوسته X مقادیر مورد انتظار برای احتمال بزرگترین عضو نمونه بدست بیاورید:

حل: می بایستی $E[F_X(X_{(n)})]$ را بدست بیاوریم داریم:

$$f_{X(n)}(t) = n(F_X(t))^{n-1} f_X(t)$$

$$E[f_X(X_{(n)})] = \int_{-\infty}^{+\infty} f_X(t) n(F_X(t))^{n-1} f_X(t) dt$$

$$= \int_{-\infty}^{+\infty} (f_X(t))^{n+1} dt = \frac{n}{n+1} - 0 = \frac{n}{n+1}$$

۹-۱۴

۴.۹. توزیع نمونه‌ای $\frac{\bar{X} - \mu}{\frac{\delta}{\sqrt{n}}}$

در این فصل نشان دادیم که توزیع $\frac{\bar{X} - \mu}{\frac{\delta}{\sqrt{n}}}$ نرمال استاندارد است، در صورتی که $X_1, X_2, \dots, X_n$ از جامعه‌ای نرمال با میانگین μ و واریانس δ^2 انتخاب شده باشد. همینطور نشان دادیم که در صورتی که تعداد نمونه‌ها زیاد باشد ($n > 30$) توزیع $\frac{\bar{X} - \mu}{\frac{\delta}{\sqrt{n}}}$ مستقل از توزیع جامعه دارای

توزیع نرمال استاندارد می باشد.

حال اگر واریانس یک جامعه مجهول باشد بناچار می بایستی از واریانس نمونه S^2 در محاسبه $\frac{\bar{X} - \mu}{\frac{\delta}{\sqrt{n}}}$ استفاده کنیم. به این ترتیب می بایستی ابتدا

تابع توزیع $\frac{\bar{X} - \mu}{\frac{\delta}{\sqrt{n}}}$ را بدست بیاوریم.

هر گاه Z یک متغیر تصادفی نرمال استاندارد باشد $Z \sim N(0, 1)$ و Y نیز یک متغیر تصادفی خی دو با n درجه آزادی باشد $(Y \sim \chi_n^2)$ و Z و Y از یکدیگر مستقل باشند در این صورت متغیر تصادفی $T = \frac{Z}{\sqrt{\frac{Y}{n}}}$ را یک متغیر تصادفی t با n درجه آزادی می‌نامیم و بصورت $T \sim t_{(n)}$ نمایش می‌دهیم.

۹-۱۵

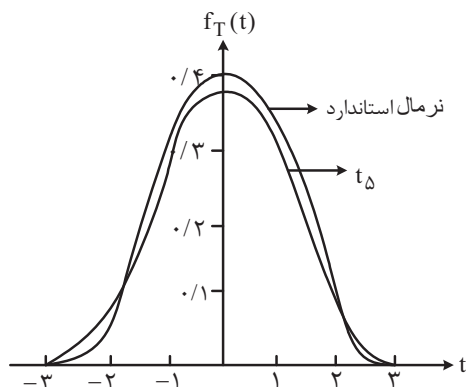
تابع چگالی متغیر تصادفی T بصورت زیر می‌باشد:

$$f_T(t) = \frac{T\left(\frac{n+1}{2}\right)}{T\left(\frac{n}{2}\right) \sqrt{n\pi}} \left[1 + \frac{t^2}{n}\right]^{-\frac{n+1}{2}} \quad -\infty < t < +\infty$$

میانگین و واریانس متغیر تصادفی t_n عبارتند از:

$$E[t_n] = 0, \quad \text{var}(t_n) = \frac{n}{n-2} \quad n > 2$$

نمودار منحنی متغیر تصادفی t_n بسیار مشابه متغیر تصادفی نرمال استاندارد می‌باشد به شکل زیر که متغیر تصادفی t_5 را با نرمال استاندارد مقایسه می‌کند توجه کنید.



در واقع با افزایش درجه آزادی متغیر تصادفی t_n نمودار منحنی هر چه بیشتر به نمودار منحنی نرمال استاندارد نزدیک می‌شود.

۹-۱۶

توجه: اگر $X \sim t_1$

توزیع t_1 بصورت زیر می‌باشد:

$$f_X(x) = \frac{1}{\pi(1+x^2)} \quad -\infty < x < +\infty$$

حال توجه کنید که معکوس توزیع X یعنی $Y = \frac{1}{X}$ برابر است با:

$$f_Y(y) = \frac{1}{\pi(1+y^2)}$$

به عبارت دقیق‌تر عکس یک متغیر t_1 باز هم t_1 می‌باشد.

حال به محاسبه توزیع $\frac{\bar{X}-\mu}{\frac{\delta}{\sqrt{n}}}$ می پردازیم. طبق تعریف می دانیم اگر $Z \sim N(0, 1)$ و $Y \sim \chi_n^2$ متغیر تصادفی $T = \frac{Z}{\sqrt{\frac{Y}{n}}}$ یک متغیر

تصادفی t_n می باشد حال در نظر بگیرید: $Z = \frac{\bar{X}-\mu}{\frac{\delta}{\sqrt{n}}}$ و $Y = \frac{(n-1) S^2}{\delta^2}$ قبلاً نشان دادیم که Z و Y به ترتیب متغیرهای تصادفی نرمال

استاندارد و χ_{n-1}^2 می باشند در نتیجه متغیر تصادفی $T = \frac{Z}{\sqrt{\frac{Y}{n-1}}}$ می بایستی یک متغیر تصادفی t با $n-1$ درجه آزادی می باشد. با ساده نمودن

عبارت T داریم:

$$T = \frac{\frac{\bar{X}-\mu}{\frac{\delta}{\sqrt{n}}}}{\sqrt{\frac{(n-1) S^2}{\delta^2 (n-1)}}} = \frac{\frac{\bar{X}-\mu}{\frac{\delta}{\sqrt{n}}}}{\sqrt{\frac{S^2}{\delta^2}}} = \frac{\bar{X}-\mu}{\frac{S}{\sqrt{n}}}$$

در نتیجه متغیر تصادفی $T = \frac{\bar{X}-\mu}{\frac{S}{\sqrt{n}}}$ یک متغیر تصادفی t با $n-1$ درجه آزادی می باشد.

توجه کنید در صورتی که نمونه های تصادفی $X_1, X_2, \dots, X_n$ از جامعه ای نرمال باشند رابطه $T = \frac{\bar{X}-\mu}{\frac{S}{\sqrt{n}}}$ یک متغیر تصادفی t_{n-1} می باشد.

اما اگر جامعه نرمال نباشد یا توزیع آنرا ندانیم یا n به اندازه کافی بزرگ ($n > 30$) باشد آنگاه باز هم $T = \frac{\bar{X}-\mu}{\frac{S}{\sqrt{n}}}$ یک متغیر تصادفی t_{n-1}

می باشد حتی در این حالت می توان با توجه به تشابه نمودار منحنی های t_{n-1} و نرمال استاندارد می توان گفت که توزیع $T = \frac{\bar{X}-\mu}{\frac{S}{\sqrt{n}}}$ برای $n > 30$

همان نرمال استاندارد می باشد.

۹-۱۷

محاسبه مقادیر احتمال t_n .

برای بدست آوردن مقادیر مختلف احتمال t_n از جدول احتمالات $F_t(x)$ استفاده می کنیم بطوریکه $F_t(x)$ برابر است با:

$$F_t(x) = \int_{-\infty}^x \frac{\Gamma(\frac{n+1}{2})}{\Gamma(\frac{n}{2}) \sqrt{n\pi}} (1+x^2)^{-\frac{(n+1)}{2}} dx = 1 - \alpha$$

$$F_t(x) = p(T < x)$$

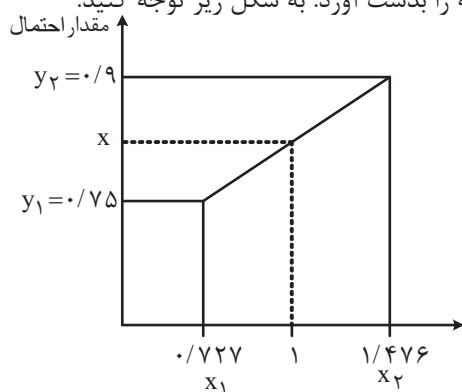
مقادیر هر یک از ردیف یا معادل با درجه آزادی و مقادیر ستون ها معادل با مقدار احتمال $F_t(x)$ می باشد. جدول زیر بخشی از جدول تابع توزیع t را نشان می دهد:

$1-\alpha$ n	۰/۶	۰/۷۵	۰/۹۰	۰/۹۵	۰/۹۷۵
۱	۰/۶۹۰	۰/۷۰۰	۰/۷۱۰	۰/۷۲۰	۰/۷۳۰

حال با توجه به جدول به عنوان مثال $p(t_2 < 2/92)$ برابر با $0/95$ می‌باشد توجه کنید که مقادیر احتمال از عناوین هر یک از ستون‌ها بدست می‌آیند. برای بدست آوردن سایر مقادیر احتمال از تقریب زدن استفاده می‌کنیم.

۹-۱۸

به عنوان مثال برای بدست آوردن مقدار احتمال $p(t_5 < 1)$ از آنجا که در جدول عدد ۱ در ردیف پنجم ما بین دو مقدار $1/476$ و $0/727$ می‌باشد می‌بایستی با استفاده از روش خطی مقدار این احتمال را تقریب بزیم در این روش ما بین دو نقطه $(0/727, 0/75)$ و $(1/476, 0/9)$ یک خط در نظر می‌گیریم، با نوشتن معادله خط به سادگی می‌توان مقادیر احتمال ما بین این دو نقطه را بدست آورد. به شکل زیر توجه کنید.



$$y = \frac{y_2 - y_1}{x_2 - x_1} (x - x_1) + y_1 \quad \text{از: معادله خط عبارتست}$$

بنابراین:

$$y = \frac{0/9 - 0/75}{1/476 - 0/727} (x - 0/727) + 0/75$$

حال برای بدست آوردن احتمال هر یک از مقادیر ما بین x_1, x_2 کافیست آن را به جای x در معادله بالا قرار دهیم تا y که مقدار احتمال تقریبی می‌باشد بدست آید.

برای احتمال $p(t_5 < 1)$ داریم:

$$p(t_5 < 1) = y = \frac{0/9 - 0/75}{1/476 - 0/727} (1 - 0/727) + 0/75 = 0/8$$

مثال: تعداد مراجعه کنندگان به یک فروشگاه در طول روز یک متغیر تصادفی نرمال می‌باشد. مدیر فروشگاه می‌داند که بطور متوسط روزانه ۱۲۰ نفر به فروشگاه مراجعه می‌کنند. تعداد مراجعه کنندگان به فروشگاه در طول یک هفته بصورت زیر بدست آمده است:

۹۰, ۱۱۰, ۷۵, ۱۳۰, ۱۵۰, ۱۲۰, ۱۰۰

مطلوبست:

الف) احتمال اینکه میانگین تعداد مراجعه کنندگان حداقل ۱۲۵ نفر در روز باشد؟

ب) احتمال اینکه تفاوت میانگین نمونه‌ها با میانگین جامعه حداکثر ۲ واحد باشد؟

حل: الف) ابتدا از روی نمونه‌های بدست آمده مقدار $\bar{X}$ و S^2 را بدست می‌آوریم. داریم:

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n x_i = \frac{1}{7} (90 + 110 + 75 + 130 + 150 + 120 + 100) = 110/7$$

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2 = \frac{1}{6} [(90 - 110/7)^2 + (75 - 110/7)^2 + (130 - 110/7)^2 + (150 - 110/7)^2 + (120 - 110/7)^2 + (100 - 110/7)^2] = 636/9 \Rightarrow S = 25/2$$

حال می‌بایستی مقدار احتمال $p(\bar{X} > 125)$ را بدست بیاوریم: (با توجه به اینکه واریانس جامعه را نمی‌دانیم بنابراین می‌بایستی از توزیع t استفاده کنیم).

$$p(\bar{X} > 110) = p(\bar{X} - \mu > 125 - 120) = p\left(\frac{\bar{X} - \mu}{\frac{S}{\sqrt{n}}} > \frac{125 - 120}{\frac{25/2}{\sqrt{7}}}\right) = p\left(\frac{\bar{X} - \mu}{\frac{S}{\sqrt{n}}} > \frac{5}{9/52}\right) = p(t_{\epsilon} > 0/525) = 1 - p(t_{\epsilon} < 0/525) = 1 - 0/68 = 0/32$$

ب) در این حالت می‌بایستی احتمال $p(\bar{X} - \mu < 2)$ را بدست بیاوریم:

$$p(\bar{X} - \mu < 2) = p\left(\frac{\bar{X} - \mu}{\frac{S}{\sqrt{n}}} < \frac{2}{9/52}\right) = p(t_{\epsilon} < 0/21) \approx 0/58$$

۹-۲۰

توزیع نمونه‌ای نسبت واریانسهای نمونه

اگر از دو جامعه نرمال دو نمونه مستقل به اندازه‌های n_1 و n_2 بگیریم برای بدست آوردن نسبت واریانسهای دو نمونه $\left(\frac{S_2}{S_1}\right)$ یا $\left(\frac{S_1}{S_2}\right)$ از توزیعی به نام F (فیشر) استفاده می‌کنیم.

تعریف: اگر U یک متغیر تصادفی خی دو با m درجه آزادی $(U \sim \chi_m^2)$ و V یک متغیر تصادفی خی دو با n درجه آزادی $(V \sim \chi_n^2)$ باشند.

در این صورت توزیع متغیر تصادفی $F = \frac{U/m}{V/n}$ را یک متغیر F با m و n درجه آزادی می‌نامند و با نماد $F \sim F_{m,n}$ نمایش می‌دهند. تابع چگالی احتمال توزیع $F_{m,n}$ بصورت زیر می‌باشد:

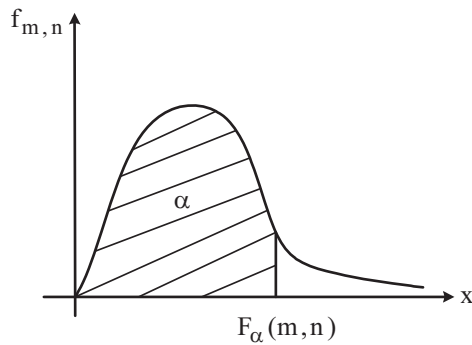
$$f_F(x) = \frac{\Gamma\left(\frac{m+n}{2}\right) \left(\frac{m}{n}\right)^{\frac{m}{2}}}{\Gamma\left(\frac{m}{2}\right) \Gamma\left(\frac{n}{2}\right)} x^{\frac{m}{2}-1} \left(1 + \frac{m}{n}x\right)^{-\frac{m+n}{2}} \quad x > 0$$

۹-۲۱

مقادیر میانگین و واریانس توزیع $F_{m,n}$ عبارتند از:

$$E[X] = \frac{n}{n-2} \quad n > 2, \quad \text{var}(X) = 2 \left(\frac{1}{n} + \frac{1}{m}\right)$$

شکل نمودار تابع توزیع $F_{m,n}$ بصورت زیر می باشد:



حال برای بدست آوردن توزیع نسبت $\frac{S_1^2}{S_2^2}$ بصورت زیر عمل می کنیم:

X و Y را دو جمعیت نرمال با واریانسهای δ_1^2 و δ_2^2 در نظر می گیریم

$$X \sim N(\mu_1, \delta_1^2) \quad ; \quad Y \sim N(\mu_2, \delta_2^2) \quad \text{فرض}$$

می دانیم:

$$V = \frac{(n_1 - 1) S_1^2}{\delta_1^2} \sim \chi_{(n_1 - 1)}^2$$

$$V = \frac{(n_2 - 1) S_2^2}{\delta_2^2} \sim \chi_{(n_2 - 1)}^2$$

حال با توجه به تعریف تابع توزیع F داریم:

$$F = \frac{\frac{U}{n-1}}{\frac{V}{n_2-1}} = \frac{\frac{(n_1-1) S_1^2}{\delta_1^2 (n_1-1)}}{\frac{(n_2-1) S_2^2}{\delta_2^2 (n_2-1)}} = \frac{\frac{S_1^2}{\delta_1^2}}{\frac{S_2^2}{\delta_2^2}} = \frac{S_1^2 \delta_2^2}{S_2^2 \delta_1^2} \sim F_{n_1-1, n_2-1}$$

۹-۲۲

مثال: اگر X یک متغیر تصادفی t با n درجه آزادی باشد مربع متغیر تصادفی X دارای چه توزیعی می باشد؟

حل: فرض می کنیم: $Z \sim N(0, 1)$ و $Y = \chi_n^2$ در این صورت داریم:

$$X = t_n \sim \frac{Z}{\sqrt{\frac{Y}{n}}}$$

$$\Rightarrow X^2 = \frac{Z^2}{\frac{Y}{n}}$$

از آنجا که مربع یک متغیر تصادفی

نرمال استاندارد یک متغیر تصادفی χ_1^2 دو با یک درجه آزادی است. بنابراین:

$$X^2 = \frac{\chi_1^2}{\frac{Y}{n}} = \frac{\chi_1^2}{\frac{\chi_n^2}{n}} \sim F_{1,n}$$

بنابراین مربع یک متغیر تصادفی t با n درجه آزادی یک متغیر تصادفی F با $n, 1$ درجه آزادی می‌باشد.

۹-۲۳

مثال: تعداد تصادفات در دو شهر A و B متغیر تصادفی نرمال به ترتیب با واریانس 20 و 22 تصادف در روز می‌باشد. اگر از دو جامعه فوق دو نمونه به ترتیب با حجم 21 و 31 نمونه انتخاب کرده باشیم و S_A^2 , S_B^2 به ترتیب واریانسهای نمونه‌های گرفته شده از شهر A و B باشد محاسبه احتمال اینکه S_A^2 حداقل $1/5$ برابر S_B^2 باشد؟

حل: داریم:

$$\sigma_A^2 = 20 \quad ; \quad \sigma_B^2 = 22$$

می‌بایستی احتمال $p(S_A^2 \geq 1/5 S_B^2)$ را محاسبه کنیم داریم:

$$p(S_A^2 \geq 1/5 S_B^2) = p\left(\frac{S_A^2}{S_B^2} \geq 1/5\right)$$

$$p\left(\frac{\frac{S_A^2}{\sigma_A^2}}{\frac{S_B^2}{\sigma_B^2}} \geq 1/5 \quad \left(\frac{22}{20}\right)\right) = p\left(\frac{S_A^2 \sigma_B^2}{S_B^2 \sigma_A^2} \geq 1/65\right)$$

$$= 1 - p(F_{20,30} < 1/65) \approx 1 - 0.09 = 0.91$$

۹-۲۴

تقریب توزیع F به توزیع خی دو.

در صورتی که در یک متغیر تصادفی $F_{m,n}$ شرط $n \rightarrow \infty$ برقرار باشد می‌توانیم برای محاسبه مقادیر احتمال mF از متغیر تصادفی χ_m^2 استفاده نمایم. به عبارت دیگر $mF_{m,n}$ با شرط $n \rightarrow \infty$ دارای توزیع χ_m^2 می‌باشد.

توزیع نمونه‌ای اختلاف میانگین‌ها

با استفاده از توزیع F احتمالات مربوط به نسبت واریانسهای دو نمونه را بدست آوریم حال می‌خواهیم برای دو نمونه گرفته شده از دو جامعه نرمال توزیع اختلاف میانگین‌ها را بدست آوریم.

از دو جامعه به ترتیب با میانگین‌های μ_1, μ_2 و واریانسهای σ_1^2, σ_2^2 دو نمونه به اندازه n_1, n_2 انتخاب می‌کنیم. در این صورت برای میانگین نمونه‌ها داریم:

$$\bar{X}_1 \sim N\left(\mu_1, \frac{\sigma_1^2}{n_1}\right)$$

$$\bar{X}_2 \sim N\left(\mu_2, \frac{\sigma_2^2}{n_2}\right)$$

از آنجا که $\bar{X}_2, \bar{X}_1$ از یکدیگر مستقل می‌باشد و $\bar{X}_1 - \bar{X}_2$ نیز یک ترکیب خطی از دو متغیر تصادفی نرمال می‌باشد بنابراین $\bar{X}_1 - \bar{X}_2$ نیز

$$\bar{X}_1 - \bar{X}_2 \sim N\left(\mu_1 - \mu_2, \frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}\right)$$

نرمال است و توزیع آن عبارتست از:

با تبدیل آن به متغیر تصادفی نرمال استاندارد داریم:

$$Z = \frac{(\bar{X}_1 - \bar{X}_2) - (\mu_1 - \mu_2)}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} \sim N(0, 1)$$

حتی اگر دو جمعیت نرمال نباشند با استفاده از قضیه حد مرکزی می‌دانیم که اگر داشته باشیم $n_1, n_2 \geq 30$ باز هم هر یک از $\bar{X}_1, \bar{X}_2$ دارای توزیع تقریبی نرمال می‌باشند. و در نتیجه $\bar{X}_1 - \bar{X}_2$ نیز دارای توزیع نرمال خواهد بود. و روابط فوق برقرار خواهند بود.

۹-۲۵

مثال: در مثال قبل اگر متوسط تعدد تصادفات در شهر A ۱۰۰ و در شهر B ۱۲۰ تصادف در روز باشند. اگر دو نمونه به حجم ۳۰ انتخاب کنیم در این صورت مطلوبست احتمال اینکه میانگین تعداد تصادفات در شهر B حداقل ۱۸ تصادف از شهر A بیشتر باشد

حل: با توجه به مفروضات مساله داریم:

$$\begin{aligned} \mu_A &= 100 & \sigma_A^2 &= 20 & n_A &= 30 \\ \mu_B &= 120 & \sigma_B^2 &= 22 & n_B &= 30 \\ \bar{X}_A - \bar{X}_B &\sim N(100 - 120, \frac{20}{30} + \frac{22}{30}) \sim N(-20, 1/4) \end{aligned}$$

حال می‌بایستی احتمال $p(\bar{X}_B \geq \bar{X}_A + 18)$ را محاسبه کنیم:

$$\begin{aligned} p(\bar{X}_B \geq \bar{X}_A + 18) &= p(\bar{X}_A - \bar{X}_B \leq -18) \\ &= p(\bar{X}_A - \bar{X}_B - (\mu_A - \mu_B)) \leq -18 - (100 - 120) \\ &= p\left(\frac{\bar{X}_A - \bar{X}_B - (\mu_A - \mu_B)}{\sqrt{\frac{\sigma_A^2}{n_A} + \frac{\sigma_B^2}{n_B}}} \leq \frac{-2}{1/18}\right) \\ &= p(Z \leq 1/69) = 0.975 \end{aligned}$$

۹-۲۶

اگر در محاسبه توزیع اختلاف میانگین نمونه‌ها واریانس دو جامعه نامعلوم اما مساوی باشد ($\sigma_1^2 = \sigma_2^2 = \sigma^2$) می‌توانیم از واریانس نمونه‌ها ($S_1^2, S_2^2, \dots$) به عنوان تخمین σ^2 استفاده کنیم. در این حالت از میانگین وزنی این دو واریانس برای برآورد σ^2 استفاده می‌کنیم:

$$S_p^2 = \frac{(n_1 - 1) S_1^2 + (n_2 - 1) S_2^2}{(n_1 + n_2 - 2)}$$

در این حالت اگر دو جامعه نرمال باشند داریم:

$$\frac{(n_1 + n_2 - 2) S_p^2}{\sigma^2} = \frac{(n_1 - 1) S_1^2}{\sigma^2} + \frac{(n_2 - 1) S_2^2}{\sigma^2} = \chi_{(n_1 + n_2 - 2)}^2$$

حال با توجه به توزیع t داریم:

$$T = \frac{(\bar{X}_1 - \bar{X}_2) - (\mu_1 - \mu_2)}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} \sim t_{(n_1 + n_2 - 2)}$$

بنابراین در محاسبه احتمال مربوط به نمونه‌های بدست آمده از دو جامعه با واریانس نامعلوم اما مساوی می‌بایستی از توزیع t با $n_1 + n_2 - 2$ درجه آزادی استفاده کنیم.

فصل دهم

S-1

۱- برآورد

در فصل قبل با نمونه‌گیری از یک جامعه با معلوم بودن تابع چگالی و پارامترهای جامعه مثل میانگین و واریانس آشنا شدید. همچنین نحوه بدست آوردن احتمالات مربوط به آماره‌های خاص مثل میانگین، نسبت واریانسها و را نشان دادیم.

در عمل بسیاری از اوقات می‌دانیم که یک جامعه مثلاً نرمال می‌باشد اما مقدار دقیق پارامترهای جامعه را که μ و σ^2 می‌باشند نمی‌دانیم، در این حالت با نمونه‌گیری از جامعه و با استنباط آماری روی نمونه‌ها می‌توانیم مقادیر پارامترهای مجهول جامعه را با دقت زیادی بدست بیاوریم. بنابراین در این فصل به بحث پیرامون روشهای برآوردیابی می‌پردازیم.

۲. ۱. ۱. برآورد نقطه‌ای

با نمونه‌گیری از جامعه دو پارامتر میانگین نمونه‌ها و واریانس نمونه‌ها را بدست می‌آوریم. این دو پارامتر نقش اصلی را در برآورد میانگین و واریانس جامعه بازی می‌کنند. بطور کلی با استفاده از میانگین و واریانس نمونه‌ها به دو طریق می‌توان پارامترهای مجهول جامعه را برآورد نمود. در روش اول که به برآورد نقطه‌ای معروف است با برابر قرار دادن میانگین و واریانس نمونه‌ها با میانگین جامعه، این دو پارامتر مجهول را بدست می‌آوریم. اما در روش دوم که به برآورد فاصله‌ای معروف است برای پارامتر مجهول جامعه مثل میانگین یک بازه با استفاده از پارامترهای معلوم نمونه بدست می‌آوریم و نشان می‌دهیم که با احتمال زیاد پارامتر مجهول جامعه در این بازه قرار دارد.

در برآورد نقطه‌ای از دو روش زیر استفاده می‌کنیم که در ادامه با آنها آشنا می‌شوید:

۱- برآورد به روش گشتاورها

۲- برآورد به روش حداکثر احتمال (در ماکزیمم).

۳. ۱. ۱. ۱۰. برآورد به روش گشتاورها

قبلاً k امین گشتاور متغیر تصادفی X را بصورت زیر معرفی نمودیم:

$$m_k = E[X^K]$$

به همین ترتیب می‌توانیم k امین گشتاور نمونه‌ها را بصورت زیر معرفی کنیم:

$$m'_k = \frac{1}{n} \sum_{i=1}^n X_i^k$$

n : تعداد نمونه‌ها

با توجه به اینکه k امین گشتاور متغیر تصادفی X به تمام مقادیر احتمال متغیر X وابسته است و همچنین k امین گشتاور نمونه‌ها نیز به هر یک از مقادیر نمونه‌ها وابسته است بنابراین این طور بنظر می‌رسد که با استفاده از گشتاورها بتوان بنحوی پارامترهای مجهول جامعه را تقریب زد. در این روش به این صورت عمل می‌کنیم که اگر p عدد پارامتر مجهول داشته باشیم به ترتیب گشتاور اول نمونه را با گشتاور اول متغیر تصادفی X ، گشتاور دوم نمونه را با گشتاور دوم متغیر تصادفی و گشتاور p ام نمونه را با گشتاور p ام متغیر تصادفی X برابر قرار می‌دهیم به این ترتیب دستگاه معادلات زیر بدست می‌آید:

$$\begin{cases} m_1 = m'_1 \\ m_2 = m'_2 \\ \vdots \\ m_k = m'_k \end{cases} \quad k=1, 2, \dots, p$$

اگر $\theta_1, \theta_2, \dots, \theta_p$ مقادیر مجهول پارامترها باشد با حل دستگاه فوق هر یک از مقادیر θ_i بدست می‌آید.

حال به مثال زیر توجه کنید:

S-2

مثال ۴: می‌دانیم تعداد مراجعه کنندگان به یک پمپ بنزین در طول روز یک متغیر تصادفی پواسون با پارامتر مجهول λ می‌باشد. اگر در طول ۱۰ روز مشاهدات زیر را برای تعداد مراجعات به پمپ بنزین بدست آورده باشیم مطابقت تعیین پارامتر مجهول λ ؟
 ۱۲۰, ۹۰, ۸۵, ۱۱۰, ۱۳۵, ۹۵, ۱۰۰, ۱۱۱, ۱۱۵, ۹۸
 حل: برای متغیر تصادفی پواسون داریم:

$$m_1 = E[X'] = E[X] = \lambda$$

از آنجا که تنها یک پارامتر مجهول λ داریم استفاده از اولین گشتاور کفایت می‌کند. حال اولین گشتاور نمونه را بدست می‌آوریم:

$$m'_1 = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X} = \frac{1}{10} (120 + 90 + 85 + 110 + 135 + 95 + 100 + 111 + 115 + 98) = 105/9$$

حال با برابر قرار دادن اولین گشتاور نمونه با اولین گشتاور متغیر تصادفی X مقدار λ بدست می‌آید:

$$m_1 = m'_1 = 105/9 \Rightarrow \lambda = 105/9$$

مثال ۵: در یک سری مسابقات تیراندازی n مرتبه به هدف شلیک می‌ند که احتمال اصابت گلوله به هدف p می‌باشد. نتایج حاصل از ۱۰ مرتبه شرکت تیرانداز در مسابقات بصورت زیر می‌باشد: (نتایج بر حسب تعداد دفعات اصابت گلوله به هدف می‌باشند)
 ۷, ۵, ۶, ۸, ۸, ۶, ۷, ۴, ۹, ۸

مطلوبست برآورد پارامترهای مجهول n و p ؟

ابتدا توجه می‌نیم که در این مثال تعداد پیروزی‌ها در n مرتبه انجام آزمایشهای مستقل برنولی با احتمال پیروزی p مورد نظر است. بنابراین متغیر تصادفی X دو جمله‌ای با پارامترهای n و p می‌باشد.

در این مثال دو پارامتر مجهول داریم بنابراین از گشتاور مرتبه اول و دوم استفاده می‌کنیم:

$$m_1 = E[X] = np$$

$$m_2 = E[X^2] = \sigma^2 + \mu^2 = npq + n^2 p^2 = n(n-1)p + np = np(np + 1 - p)$$

توجه کنید که مقادیر اول و دوم گشتاور متغیر تصادفی X همواره با استفاده از میانگین و واریانس متغیر تصادفی X بدست می‌آیند. به عبارتی در برآورد به روش گشتاورها مستقل از اینکه توزیع متغیر تصادفی X چه باشد همواره داریم:

$$\begin{cases} m_1 = \mu \\ m'_1 = \bar{X} \end{cases} \Rightarrow \mu = \bar{X}$$

$$\begin{cases} m_2 = \sigma^2 + \mu^2 \\ m'_2 = \frac{1}{n} \sum_{i=1}^n \bar{X}_i^2 \end{cases} \Rightarrow \sigma^2 = \frac{1}{n} \sum_{i=1}^n (\bar{X}_i - \bar{X})^2 = \frac{1}{n} \sum_{i=1}^n \bar{X}_i^2 - \bar{X}^2$$

یعنی برآورد گشتاوری میانگین و واریانس هر توزیعی برابر است با میانگین و واریانس نمونه‌ها. به همین دلیل $\bar{X}$, $\frac{1}{n} \sum (X_i - \bar{X})^2$ برآوردهای توزیع - آزاد نامیده می‌شوند، یعنی برآوردهایی که مستقل از توزیع تصادفی X می‌باشند.

۶ حال مقادیر گشتاور اول و دوم نمونه را بدست می‌آوریم:

$$m'_1 = \frac{1}{n} \sum_{i=1}^n X_i = \frac{1}{10} (7 + 5 + 6 + 8 + 8 + 6 + 7 + 4 + 9 + 8) = 6/8$$

$$m'_2 = \frac{1}{n} \sum_{i=1}^n X_i^2 = \frac{1}{10} (49 + 25 + 36 + 64 + 64 + 36 + 49 + 16 + 81 + 64) = 48/4$$

به این ترتیب دستگاه معادلات زیر بدست می‌آید:

$$\begin{cases} m_1 = m'_1 & \Rightarrow n p = 6/8 \\ m_2 = m'_2 & \Rightarrow n p = (n p + 1 - p) = 48/4 \end{cases}$$

با حل دستگاه بدست می‌آوریم:

$$6/8 (6/8 + 1 + p) = 48/4 \Rightarrow 53/0.4 - 6/8 p = 48/4 \Rightarrow p = 0.68$$

$$n p = 6/8 \rightarrow n (0.68) = 6/8 \Rightarrow n = 10$$

بنابراین تیرانداز فوق در هر مسابقه می‌بایستی حدوداً ۱۰ مرتبه به هدف شلیک کرده باشد که احتمال موفقیت وی در هر شلیک ۰/۶۸ می‌باشد.

برآورد به روش گشتاورها همواره نتایج دقیق و رضایت بخشی نمی‌دهد به مثال بعد در این زمینه توجه کنید:

۷ مثال ۳: فرض کنید متغیر تصادفی X در بازه $[0, a]$ بصورت یکنواخت توزیع شده باشد در این صورت مقدار a را به ازای نمونه‌های بدست آمده

زیر برآورد کنید:

الف) ۵ و ۴ و ۳ و ۱ و ۸ و ۵ و ۴ و ۲ و ۰ و ۱

ب) ۵ و ۱۰ و ۴۰

حل: برای متغیر تصادفی پیوسته یکنواخت داریم:

$$m_1 = E[X] = \int_0^a x \cdot \frac{1}{a} dx = \frac{1}{a} \left(\frac{x^2}{2} \right) \Big|_0^a = \frac{a}{2}$$

همینطور برای نمونه‌ها داریم:

$$m'_1 = \frac{1}{n} \sum_{i=1}^n X_i = \frac{1}{10} (1+0+2+4+5+8+1+3+4+5) = 3/3$$

$$a = 2\bar{X} \quad \Leftarrow \quad m'_1 = \bar{X} = \frac{a}{2}$$

بنابراین در حالت کلی:

یعنی برآورد a عبارتست از دو برابر میانگین نمونه‌های بدست آمده که در اینجا با توجه به نمونه‌های الف a برابر می‌شود با: $a = 2\bar{X} = 2 \times 3/3 = 6/6$.

حال توجه کنید که $a = 6/6$ یعنی نمونه‌ها در اصل از بازه $[0, 6/6]$ انتخاب شده‌اند در حالی که در بین نمونه‌ها عدد ۸ موجود می‌باشد که

داخل بازه فوق نمی‌باشد و این یعنی برآورد به روش گشتاورها دارای کمبودهایی می‌باشد. همین حالت برای نمونه‌های (ب) نیز صادق است:

(ب)

$$\bar{X} = \frac{1}{3} (40+10+5) = 18/3 \Rightarrow a = 2\bar{X} = 36/6$$

بنابراین برآورد به روش گشتاورها در بعضی موارد نتایج دلخواه را بدست نمی‌دهد در نتیجه می‌بایستی از معیاری استفاده کنیم که میزان کارایی یک

روش برآورد را نشان دهد تا بتوان میان روشهای مختلف برآورد، بهترین روش را برای مسایل مختلف انتخاب نمود.

S-۳

۲.۱.۱.۱. برآورد به روش حداکثر احتمال

متغیر تصادفی X با تابع چگالی احتمال $f_X(x)$ و پارامترهای مجهول $\theta_1, \theta_2, \dots, \theta_p$ را در نظر بگیرید. از این متغیر تعداد n نمونه استخراج

می‌کنیم که عبارتند از $x_1, x_2, \dots, x_n$ حال $X_1, X_2, \dots, X_n$ را متغیرهای تصادفی متناظر با نمونه‌ها در نظر می‌گیریم.

می‌دانیم تابع چگالی احتمال توام متغیرهای تصادفی $X_1, \dots, X_n$ مقادیر احتمال وقوع نمونه‌ها را بدست می‌دهد. مقادیری از θ_i ها که تابع چگالی

احتمال توام X_i ها را ماکزیمم می‌کند برآورد پارامترهای مجهول θ_i می‌باشد. زیرا به ازای ماکزیمم شدن تابع چگالی احتمال توام X_i ها در واقع

احتمال وقوع نمونه‌ها به بیشترین مقدار خود می‌رسد. برای بدست آوردن برآورد θ_i ها فرض می‌کنیم $L(\theta_1, \theta_2, \dots, \theta_p)$ تابع چگالی احتمال

توام متغیرهای تصادفی X_i باشد در این صورت با توجه به اینکه X_i ها از یکدیگر مستقل باشند تابع چگالی احتمال توام آنها عبارتست از:

$$L(\theta_1, \theta_2, \dots, \theta_p) = f_{X_1}(x_1) \cdot f_{X_2}(x_2) \cdots f_{X_n}(x_n)$$

حال با مشتق‌گیری از $L(\theta_1, \dots, \theta_p)$ نسبت به θ_i ها و برابر صفر قرار دادن این مشتق‌ها دستگاهی از معادلات بدست می‌آید که با حل این دستگاه مقادیر هر یک از θ_i ها بدست می‌آید. معمولاً در محاسبات برای سادگی از $\ln(L)$ مشتق می‌گیریم. به مثال زیر توجه کنید:

مثال ۹: نمونه‌های $X_1, X_2, \dots, X_n$ را از متغیر تصادفی برنولی X داریم مطلوبست: برآورد پارامتر p به روش حداکثر احتمال (MLE)؟

حل: تابع چگالی متغیر تصادفی برنولی X به صورت زیر می‌باشد:

$$f_X(x) = p^x (1-p)^{1-x} \quad x = 0, 1$$

سایر مقادیر

به این ترتیب تابع چگالی احتمال توام نمونه‌ها عبارتست از:

$$L(p) = f_{X_1}(x_1) \cdot f_{X_2}(x_2) \cdots f_{X_n}(x_n) = \prod_{i=1}^n f_X(x_i) = \prod_{i=1}^n p^{x_i} (1-p)^{1-x_i}$$

$$p^{\sum_{i=1}^n x_i} (1-p)^{\sum_{i=1}^n (1-x_i)} \Rightarrow L(p) = p^{\sum_{i=1}^n x_i} (1-p)^{n - \sum_{i=1}^n x_i}$$

حال برای اینکه محاسبات ساده‌تر شوند از طرفین رابطه بالا $\ln$ می‌گیریم:

$$\begin{aligned} \ln(L(p)) &= \ln(p^{\sum x_i} (1-p)^{n - \sum x_i}) \\ &= \sum x_i \ln(p) + (n - \sum x_i) \ln(1-p) \end{aligned}$$

فرض می‌کنیم $H(p) = \ln(L(p))$ در این صورت می‌بایستی معادله $\frac{dH}{dp} = 0$ را حل کنیم تا مقدار مجهول p بدست آید. بنابراین:

$$\frac{dH}{dp} = 0 \Rightarrow \sum x_i \left(\frac{1}{p}\right) - (n - \sum x_i) \left(\frac{1}{1-p}\right) = 0$$

همانطور که ملاحظه می‌کنید با بکار بردن $\ln$ در محاسبات، مشتق‌گیری بسیار ساده‌تر می‌شود. حال داریم:

$$\begin{aligned} \sum x_i \left(\frac{1}{p}\right) &= (n - \sum x_i) \left(\frac{1}{1-p}\right) \\ \Rightarrow \frac{1-p}{p} &= \frac{n - \sum x_i}{\sum x_i} \Rightarrow \frac{1}{p} - 1 = \frac{n}{\sum x_i} - 1 \\ \Rightarrow \frac{1}{p} &= \frac{n}{\sum x_i} \Rightarrow \hat{p} = \frac{1}{n} \sum x_i \end{aligned}$$

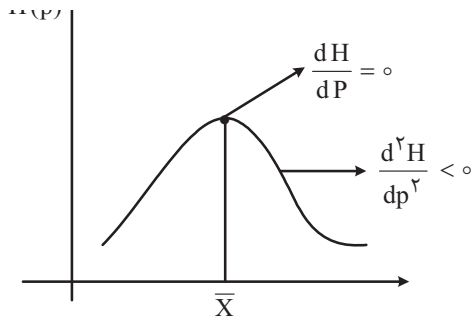
($\hat{p}$ برآورد پارامتر مجهول p می‌باشد).

بنابراین برآورد پارامتر مجهول p برابر است با $\bar{X} = \frac{1}{n} \sum x_i$ توجه کنید مقدار $\hat{p} = \bar{X}$ بدست آمده تنها در صورتی جواب درستی می‌باشد که نقطه

$\bar{X}$ واقعاً مقدار ماکزیمم تابع $H(p)$ باشد. مثلاً ممکن است این نقطه، نقطه مینیمم تابع $H(p)$ باشد. برای اطمینان از اینکه $\hat{p} = \bar{X}$ مقدار ماکزیمم $H(p)$ می‌باشد باید شرط زیر برقرار باشد:

$$\frac{d^2 H}{dp^2} < 0$$

یعنی تقعر تابع $H(p)$ حول $\hat{p} = \bar{X}$ روبه پایین باشد به این ترتیب $\hat{p} = \bar{X}$ نقطه ماکزیمم تابع می‌باشد. در این رابطه به شکل زیر توجه کنید:



حال داریم:

$$\frac{d^2H}{dp^2} = \sum x_i \left(\frac{-1}{p^2} \right) - (n - \sum x_i) \left(\frac{1}{(1-p)^2} \right)$$

که به ازای هر مقدار p منفی است. بنابراین $\bar{X}$ به عنوان برآورد p در متغیر تصادفی برنولی می‌باشد.

۱۱ حال برآورد p را با استفاده از روش گشتاورها بدست می‌آوریم و نتیجه را با روش حداکثر احتمال مقایسه می‌کنیم:

مطابق روش گشتاورها داریم:

$$m_1 = p$$

$$\Rightarrow m_1 = m'_1 \Rightarrow p = \bar{X}$$

$$m'_1 = \frac{1}{n} \sum x_i$$

همانطور که ملاحظه می‌کنید نتیجه برآورد از هر دو روش یک نتیجه را بدست می‌دهد. بطور کلی از آنجا که روش حداکثر معمولاً برآوردهای بهتری می‌دهد، هرگاه نتیجه دو روش برآورد با یکدیگر برابر نبود، نتیجه روش برآورد حداکثر احتمال را ملاک قرار می‌دهیم. در برآورد به روش گشتاورها نشان دادیم که این روش برای مواردی مثل مثال ۳ برآورد مطلوبی ارایه نمی‌کند بنابراین در مثال بعدی مثال ۳ را با روش برآورد حداکثر احتمال حل می‌کنیم:

S ۴

۱۲ مثال ۵: متغیر X یک متغیر تصادفی یکنواخت در بازه $(0, a)$ می‌باشد. مطلوبست برآورد پارامتر مجهول a به روش حداکثر احتمال؟

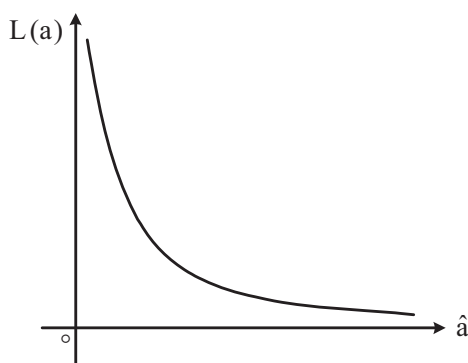
حل: تابع چگالی احتمال X عبارتست از:

$$f_X(x) = \frac{1}{a} \quad \text{برای } 0 < x < a$$

با در نظر گرفتن یک نمونه تصادفی به اندازه n تابع چگالی احتمال توأم بصورت زیر بدست می‌آید:

$$L(a) = \frac{1}{a} \cdot \frac{1}{a} \cdots \frac{1}{a} = \frac{1}{a^n}$$

منحنی نمایش $L(a)$ را در شکل زیر مشاهده می‌کنید.



از روی منحنی پیداست که شیب منحنی در هیچ نقطه‌ای صفر نمی‌شود بنابراین مقدار مشتق $L(a)$ در هیچ نقطه‌ای صفر نمی‌شود. اما با توجه به نمودار می‌بینیم که هر قدر a به صفر نزدیک می‌شود مقدار $L(a)$ بیشتر می‌شود، و به عبارتی $L(a)$ به ازای کوچکترین مقدار a ماکزیمم می‌شود، از طرفی کوچکترین مقدار a نباید از بزرگترین نمونه بدست آمده کمتر باشد بنابراین برآورد a به روش حداکثر احتمال در این مثال برابر است با بزرگترین نمونه بدست آمده یعنی $\hat{a} = X_{(n)}$ امین آماره مرتب.

بوضوح این برآورد کاراتر از برآورد به روش گشتاورها می‌باشد.

۱۳ مثال ۶: اگر متغیر تصادفی X یک متغیر تصادفی پواسون با پارامتر λ باشد مطلوبست برآورد λ به روش حداکثر احتمال و به روش گشتاورها؟
حل: تابع چگالی متغیر تصادفی پواسون عبارتست از:

$$f_X(x) = \frac{e^{-\lambda} \lambda^x}{x!} \quad x = 0, 1, 2, \dots$$

$= 0$ سایر مقادیر

$$L(\lambda) = \frac{e^{-\lambda} \lambda^{x_1}}{x_1!} \cdot \frac{e^{-\lambda} \lambda^{x_2}}{x_2!} \dots \frac{e^{-\lambda} \lambda^{x_n}}{x_n!} = \frac{e^{-n\lambda} \lambda^{\sum x_i}}{x_1! x_2! \dots x_n!}$$

$$H(\lambda) = \ln(L(\lambda)) = -n\lambda + \sum x_i \ln(\lambda) - \ln(x_1! x_2! \dots x_n!)$$

$$\frac{dH}{d\lambda} = -n + \sum x_i \left(\frac{1}{\lambda}\right) = 0 \Rightarrow \hat{\lambda} = \frac{1}{n} \sum x_i = \bar{X}$$

حال این مساله را به روش گشتاورها حل می‌کنیم:

$$m_1 = \mu = \lambda$$

$$\Rightarrow m_1 = m'_1 \Rightarrow \hat{\lambda} = \bar{X}$$

$$m'_1 = \frac{1}{n} \sum x_i = \bar{X}$$

ملاحظه می‌کنید که نتایج هر دو روش در این حالت نیز یکسان می‌باشند.

۱۴ مثال ۷: اگر متغیر تصادفی X یک متغیر تصادفی نرمال با میانگین μ و واریانس σ^2 باشد مطلوبست برآورد پارامترهای μ و σ^2 به روشهای حداکثر احتمال و گشتاورها؟
حل: تابع چگالی احتمال متغیر تصادفی X عبارتست از:

$$f_X(x) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x-\mu)^2}{2\sigma^2}}$$

$$L(\mu, \sigma^2) = \prod_{i=1}^n \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{(x_i-\mu)^2}{2\sigma^2}} = \frac{1}{(\sigma\sqrt{2\pi})^n} e^{-\frac{1}{2\sigma^2} \sum_{i=1}^n (x_i-\mu)^2}$$

$$H(\mu, \sigma^2) = \ln(L(\mu, \sigma^2)) = \ln\left(\frac{1}{(\sigma\sqrt{2\pi})^n}\right) + \ln\left(e^{-\frac{1}{2\sigma^2} \sum_{i=1}^n (x_i-\mu)^2}\right)$$

$$= -\frac{n}{2} (\ln(2\pi\sigma^2)) - \frac{1}{2\sigma^2} \sum_{i=1}^n (x_i-\mu)^2$$

$$\Rightarrow H(\mu, \sigma^2) = -\frac{n}{2} (n(2\pi) - \frac{n}{2} (\sigma^2 - \frac{1}{2} \sum_{i=1}^n \frac{(x_i - \mu)^2}{\sigma^2}))$$

$$\frac{\partial H}{\partial \mu} = -\frac{1}{2} (-2) = \sum_{i=1}^n \frac{(x_i - \mu)^2}{\sigma^2} = \frac{1}{\sigma^2} \sum_{i=1}^n (x_i - \mu)^2$$

$$\begin{aligned} \frac{\partial H}{\partial \sigma^2} &= -\frac{n}{2} \left(\frac{1}{\sigma^2} \right) - \frac{1}{2} \sum_{i=1}^n (x_i - \mu)^2 \left(\frac{-1}{(\sigma^2)^2} \right) \\ &= \frac{-n}{2\sigma^2} + \frac{1}{2(\sigma^2)^2} \sum_{i=1}^n (x_i - \mu)^2 \end{aligned}$$

حال می‌بایستی معادلات زیر را حل کنیم:

$$\begin{cases} \frac{1}{\hat{\sigma}^2} \sum (x_i - \hat{\mu}) = 0 \\ -\frac{n}{2\hat{\sigma}^2} + \frac{1}{2(\hat{\sigma}^2)^2} \sum (x_i - \hat{\mu})^2 = 0 \end{cases}$$

۱۵ جواب معادله اول عبارتست از:

$$\begin{aligned} \sum (x_i - \hat{\mu}) = 0 &\Rightarrow \sum x_i - \sum \hat{\mu} = 0 \rightarrow \sum x_i - n\hat{\mu} = 0 \\ \Rightarrow \hat{\mu} &= \frac{1}{n} \sum x_i = \bar{X} \end{aligned}$$

با قرار دادن $\hat{\mu} = \bar{X}$ در معادله دوم داریم:

$$-\frac{n}{2\hat{\sigma}^2} + \frac{1}{2(\hat{\sigma}^2)^2} \sum (x_i - \bar{X})^2 = 0$$

$$n\hat{\sigma}^2 + \sum (x_i - \bar{X})^2 = 0 \Rightarrow \hat{\sigma}^2 = \frac{1}{n} \sum (x_i - \bar{X})^2$$

بنابراین برآورد پارامترهای مجهول μ و σ^2 به روش حداکثر احتمال عبارتند از:

$$\hat{\mu} = \bar{X}, \quad \hat{\sigma}^2 = \frac{n-1}{n} S^2$$

۱۶ حال برآورد را به روش گشتاورها بدست می‌آوریم:

$$m_1 = \mu$$

$$m_2 = \sigma^2 + \mu^2$$

$$m'_1 = \frac{1}{n} \sum_{i=1}^n x_i = \bar{X}$$

$$m'_2 = \frac{1}{n} \sum_{i=1}^n x_i^2$$

$$\begin{cases} m_1 = m'_1 \\ m_2 = m'_2 \end{cases} \Rightarrow \begin{cases} \mu = \bar{X} \\ \frac{1}{n} \sum_{i=1}^n x_i^2 = \sigma^2 + \mu^2 \end{cases}$$

از معادله اول داریم $\hat{\mu} = \bar{X}$ با جاگذاری در معادله دوم بدست می‌آید:

$$\frac{1}{n} \sum_{i=1}^n x_i^2 = \hat{\sigma}^2 + \bar{X}^2 \Rightarrow \hat{\sigma}^2 = \frac{1}{n} \sum_{i=1}^n x_i^2 - \bar{X}^2$$

$$\hat{\sigma}^2 = \frac{1}{n} \frac{1}{n} \sum_{i=1}^n (x_i - \bar{X})^2 = \frac{n-1}{n} S^2$$

در نتیجه برآورد پارامترهای μ و σ^2 به روش گشتاورها نیز همان نتایج برآورد به روش حداکثر را بدست می‌دهد.

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2$$

S-5

۱۶ مثال ۸: اگر متغیر تصادفی X بصورت یکنواخت در بازه $[-b, b]$ توزیع شده باشد مطلوبست برآورد b به روش گشتاورها و حداکثر احتمال؟
حل: میانگین متغیر تصادفی X عبارتست از:

$$E[X] = \int_{-b}^b dx = \frac{1}{2b} \left(\frac{x^2}{2} \right) \Big|_{-b}^b = 0$$

حال طبق روش گشتاورها می‌بایستی از حل معادله $m_1 = m'_1$ پارامتر b بدست آید اما می‌بینیم که:

$$m_1 = E[X] = 0$$

$$m'_1 = \frac{1}{n} \sum_{i=1}^n x_i = \bar{X}$$

$$m_1 = m'_1 \Rightarrow \bar{X} = 0$$

همانطور که ملاحظه می‌کنید $\bar{X} = 0$ کمکی در یافتن پارامتر مجهول به ما نمی‌کند. در واقع این مثال حالت خاصی است که در آن جهت تعیین پارامتر مجهول می‌بایستی از گشتاور دوم نمونه و متغیر تصادفی X استفاده کنیم:

$$m_2 = E[X^2] = \int_{-b}^b x^2 \frac{1}{2b} dx = \frac{1}{2b} \left(\frac{x^3}{3} \right) \Big|_{-b}^b$$

$$= \frac{1}{2b} \left(\frac{b^3}{3} + \frac{b^3}{3} \right) = \frac{b^2}{3}$$

به همین ترتیب داریم:

$$m'_2 = \frac{1}{n} \sum_{i=1}^n x_i^2 \Rightarrow m_2 = m'_2 \Rightarrow \frac{b^2}{3} = \frac{1}{n} \sum_{i=1}^n x_i^2$$

$$\Rightarrow b^2 = 3 \left(\frac{1}{n} \sum_{i=1}^n x_i^2 \right) \rightarrow \hat{b} = \sqrt{3 \left(\frac{1}{n} \sum_{i=1}^n x_i^2 \right)} = \sqrt{3 m'_2}$$

بنابراین برآورد پارامتر مجهول b در این مثال برابر $\sqrt{3 m'_2}$ می‌باشد.

۱۷ روش برآورد پارامتر b به روش حداکثر احتمال کاملاً مشابه مثال ۵ می‌باشد به این صورت که:

$$f_X(x) = \frac{1}{2b} \quad -b < x < b$$

$$L(b) = \frac{1}{2b} \cdot \frac{1}{2b} \cdots \frac{1}{2b} = \frac{1}{(2b)^n}$$

مجدداً از آنجا که $L(b)$ با نزدیک شدن به صفر ماکزیمم می‌شود بنابراین کوچکترین مقداری که b می‌تواند بخود بگیرد برآورد آن می‌باشد و از آنجا که b نمی‌تواند از بزرگترین مقدار نمونه کمتر باشد بنابراین در این مثال نیز برآورد b عبارتست از $\max |X_i|$ یا همان n امین آماره مرتب $X_{(n)}$.

۱۸. ۱. ۲. خواص برآورد کننده‌ها

در بخش روشهای گشتاورها و حداکثر احتمال را برای برآورد پارامترهای مجهول جامعه معرفی نمودیم و نشان دادیم که این دو روش همواره برآورد کننده‌های مشابه بدست نمی‌دهند، بنابراین نیازمند یک معیاری برای سنجش میزان کارایی یک برآوردگر نسبت به برآوردگر دیگری می‌باشیم، در این بخش به معرفی خواص برآورد کننده‌ها و روشهایی برای سنجش میزان کارایی آنها می‌پردازیم.

۱۹. ۱. ۲. متوسط مربع خطا (MSE)

فرض کنید $\hat{\theta}$ برآورد پارامتر θ باشد در این صورت می‌دانیم که $|\hat{\theta} - \theta|$ مقدار خطای برآورد می‌باشد. معمولاً برای تخمین میزان خطا از مربع $(\hat{\theta} - \theta)$ استفاده می‌شود و از آنجا که در اینجا n نمونه استخراج می‌شود برای هر بار استخراج این n نمونه مقدار $(\hat{\theta} - \theta)^2$ محاسبه می‌شود و میانگین آن ملاک خواهد بود. به این ترتیب معیار متوسط مربع خطا بصورت زیر بدست می‌آید:

$$MSE(\hat{\theta}) = E[(\hat{\theta} - \theta)^2]$$

توجه کنید که همواره در عمل با توجه به استخراج n نمونه و برآورد پارامتر θ مقداری خطا در برآورد پارامتر θ وجود دارد، بنابراین برای اینکه بهترین برآورد را داشته باشیم می‌بایستی خطای برآورد را حداقل کنیم. بوضوح اگر $\hat{\theta} = \theta$ باشد خطای برآورد صفر می‌باشد و بهترین حالت بدست آمده است.

برای اینکه بتوان خطای برآورد را حداقل نمود ابتدا مقدار متوسط مربع خطا را ساده می‌کنیم:

$$\begin{aligned} MSE &= E[(\hat{\theta} - \theta)^2] = E[(\hat{\theta}^2 - 2\theta\hat{\theta} - \theta)^2] \\ &= E[\hat{\theta}^2] - 2E[\theta\hat{\theta}] + E[\theta^2] \end{aligned}$$

حال یک $E^2[\hat{\theta}]$ به سمت راست عبارت بالا اضافه و کم می‌کنیم:

$$\Rightarrow MSE = E[\hat{\theta}^2] - E^2[\hat{\theta}] + E^2[\hat{\theta}] - 2E[\theta\hat{\theta}] + E[\theta^2]$$

توجه کنید که عبارت A همان واریانس $\hat{\theta}$ می‌باشد. می‌دانیم $\hat{\theta}$ خود یک آماره می‌باشد و از آنجا که آماره‌ها خود یک متغیر تصادفی می‌باشند، بنابراین واریانس $\hat{\theta}$ معنی‌دار می‌باشد.

$$\Rightarrow MSE = \text{var}(\hat{\theta}) + (E[\hat{\theta}] - \theta)^2$$

بنابراین متوسط مربع خطا بصورت فوق خلاصه می‌شود که در آن عبارت $E[\hat{\theta}] - \theta$ در واقع اندازه اختلاف مرکز توزیع $\hat{\theta}$ از θ را نشان می‌دهد.

۲۰. برای اینکه MSE حداقل شود می‌بایستی مقادیر $\text{var}(\hat{\theta})$ ، $(E[\hat{\theta}] - \theta)^2$ حداقل شوند. از طرفی می‌دانیم $\text{var}(\hat{\theta})$ ، $(E[\hat{\theta}] - \theta)^2$ هر دو عبارتی همواره مثبت می‌باشند بنابراین MSE در صورتی حداقل می‌شود که این دو عبارت حداقل شوند. حال کمترین مقداری را که این دو عبارت بخود می‌پذیرد را محاسبه می‌کنیم:

$$\text{var}(\hat{\theta}) = E[\hat{\theta}^2] - E^2[\hat{\theta}]$$

اگر $E^2[\hat{\theta}] = E[\hat{\theta}^2]$ واریانس $\hat{\theta}$ برابر صفر می‌شود که کمترین مقداری است که یک عبارت مثبت بخود می‌پذیرد. بنابراین $\min(\text{var}(\hat{\theta})) = 0$ همینطور حداقل $(E[\hat{\theta}] - \theta)^2$ با توجه به مثبت بودن عبارت زمانی رخ می‌دهد که کل عبارت صفر شود یعنی:

$$\min [(E[\hat{\theta}] - \theta)^2] = 0 \Rightarrow (E[\hat{\theta}] - \theta)^2 = 0 \Rightarrow E[\hat{\theta}] = \theta$$

می‌دانیم در عمل $\text{var}(\hat{\theta})$ هرگز صفر نمی‌شود (مگر اینکه تمام مقادیر نمونه‌ها برابر باشند که آن هم بی‌معنی است). اما رخ دادن تساوی

$E[\hat{\theta}] = \theta$ امکان‌پذیر است بنابراین از میان برآوردهای مختلف برای θ برآوردگری مناسب‌تر است که شرایط زیر را داشته باشد.

۱- واریانس برآوردگر $(\hat{\theta})$ کمترین مقدار ممکن را داشته باشد.

۲- $E[\hat{\theta}] = \theta$ باشد.

شرط دوم به برآورد کننده نااریب معروف است که به صورت زیر تعریف می‌شود:

برآورد کننده ناریب: برآورد کننده $\hat{\theta}$ برای θ ناریب نامیده می شود اگر $E[\hat{\theta}] = \theta$.

برابر تساوی $E[\hat{\theta}] = \theta$ به این معنی است که اگر مثلاً k بار از جامعه تعداد n نمونه استخراج کنیم و به ازای هر بار استخراج n نمونه مقدار $\hat{\theta}$ را محاسبه کنیم و در نهایت میانگین $\hat{\theta}_i$ ها $(E[\hat{\theta}] = \frac{1}{k} \sum_{i=1}^k \hat{\theta}_i)$ می بایستی برابر با پارامتر مجهول θ شود. برای روشن شدن مطلب به مثال زیر توجه کنید.

۲۱ مثال ۹: آیا برآوردگر بدست آمده برای توزیع نرمال، برنولی، پواسون ناریب می باشد؟

حل: قبلاً نشان دادیم که برآوردگر μ و σ^2 در توزیع نرمال با استفاده از روشهای گشتاورها و حداکثر احتمال عبارتست از:

$$\hat{\mu} = \bar{X} \quad \hat{\sigma}^2 = \frac{n-1}{n} S^2$$

برای $\hat{\mu}$ داریم:

$$E[\hat{\mu}] = E[\bar{X}] = \mu$$

توجه کنید که قبلاً نشان دادیم برای متغیر تصادفی X با هر توزیعی $E[\bar{X}]$ برابر μ می باشد. بنابراین $\hat{\mu} = \bar{X}$ در توزیع نرمال یک برآورد ناریب برای μ می باشد.

حال برای $\hat{\sigma}^2$ داریم:

$$E[\hat{\sigma}^2] = E\left[\frac{n-1}{n} S^2\right] = \frac{n-1}{n} E[S^2]$$

$$E[S^2] = \sigma^2 \quad \text{۲۲ در فصل قبل نشان دادیم که اگر } S^2 = \frac{1}{n-1} \sum_{i=1}^n (X_i - \bar{X})^2 \text{ آنگاه داریم:}$$

$$E[\hat{\sigma}^2] = \frac{n-1}{n} \sigma^2$$

بنابراین:

ملاحظه می کنید که برای $\hat{\sigma}^2$ مقدار $E[\hat{\sigma}^2]$ برابر با σ^2 نمی باشد بنابراین برآوردگر $\hat{\sigma}^2 - \frac{n-1}{n} S^2$ برآوردگر ناریبی برای σ^2 نمی باشد. در این حالت می گوییم برآوردگر $\hat{\sigma}^2$ یک برآوردگر اریب برای σ^2 می باشد.

توجه کنید که $\hat{\sigma}^2$ تنها به دلیل وجود ضریب $(\frac{n-1}{n})$ اریب می باشد، به عبارتی تعداد نمونه ها عامل اصلی اریب بودن این برآوردگر می باشد بنابراین

$$\lim_{n \rightarrow \infty} \left(\frac{n-1}{n}\right) \sigma^2 = \sigma^2$$

اگر تعداد نمونه ها زیاد باشد $(n \rightarrow \infty)$ داریم:

یعنی با شرط $n \rightarrow \infty$ برآوردگر $\hat{\sigma}^2$ یک برآوردگر ناریب برای σ^2 می باشد. به این نوع برآوردگرها، برآوردگر ناریب مجانبی می گویند که بصورت زیر تعریف می شود:

برآوردگر ناریب مجانبی: اگر $\hat{\theta}$ یک برآوردگر اریب برای θ بر اساس نمونه تصادفی n تایی باشد، می گوییم $\hat{\theta}$ یک برآوردگر ناریب مجانبی برای θ است هرگاه داشته باشیم:

$$\lim_{n \rightarrow \infty} E[\hat{\theta}] = \theta$$

۲۳ برای توزیع برنولی قبلاً بدست آوردیم که:

$$\hat{p} = \bar{X} \Rightarrow E[\hat{p}] = E[\bar{X}] = \mu = p$$

همینطور برای توزیع پواسون داریم:

$$\hat{\lambda} = \bar{X} \\ E[\hat{\lambda}] = E[\bar{X}] = \mu = \lambda$$

یک نتیجه کلی که از برآوردگر ناریب بدست می آوریم این است که هرگاه پارامتر مجهول توزیع برابر با میانگین توزیع متغیر تصادفی X باشد

(یعنی $\mu = E[X] = \theta$) همینطور از آنجا که برای هر توزیعی مستقل از نوع توزیع داریم $E[\bar{X}] = \mu$ در این صورت $\bar{X}$ همواره به عنوان یک برآوردگر نااریب این گونه از توزیعها، منظور می‌شود.

خاصیت نااریبی اگر چه یک معیار برای محک بهتر بودن یک برآوردگر می‌باشد اما از آنجا که برای یک پارامتر می‌توان چندین برآوردگر نااریب بدست آورد، بنابراین می‌بایستی معیارهای دیگری وجود داشته باشند تا بتوان با کمک آنها از میان برآوردگرهای نااریب بهترین را انتخاب نمود. برای این منظور معیار کارایی یا موثرتر بودن را معرفی می‌کنیم:

تعریف: اگر $\hat{\theta}_1, \hat{\theta}_2$ هر دو برآوردگرهایی نااریب برای پارامتر θ باشند در این صورت می‌گوییم $\hat{\theta}_1$ کاراتر یا موثرتر از $\hat{\theta}_2$ می‌باشد هرگاه:

$$\text{var}(\hat{\theta}_1) < \text{var}(\hat{\theta}_2)$$

علت استفاده از این معیار با توجه به تعریف متوسط مربع خطا (MSE) کاملاً قابل توجیه است.

مثال ۲۴: ۱۰: از یک جامعه نرمال با میانگین μ و واریانس σ^2 یک نمونه به حجم n انتخاب شده است. در این صورت کدامیک از برآوردگرهای نااریب زیر معیار مناسب‌تری برای برآورد پارامتر μ می‌باشد؟

الف) $\hat{\mu}_1 = \bar{X}$

ب) $\hat{\mu}_2 = \frac{1}{3} (X_1 + X_3 + X_5)$

حل: با توجه به اینکه هر دو برآوردگر نااریب می‌باشند بنابراین واریانس آنها را محاسبه نموده و با یکدیگر مقایسه می‌کنیم:

$$\text{var}(\hat{\mu}_1) = \text{var}(\bar{X}) = \frac{\sigma^2}{n}$$

$$\text{var}(\hat{\mu}_2) = \text{var}\left(\frac{1}{3} (X_1 + X_3 + X_5)\right) = \frac{1}{9} (\text{var}(X_1) + \text{var}(X_3) + \text{var}(X_5))$$

$$= \frac{1}{9} (3\sigma^2) = \frac{\sigma^2}{3}$$

با مقایسه دو واریانس بدست آمده داریم:

$$\frac{\sigma^2}{n} < \frac{\sigma^2}{3} \quad \text{برای } n > 3$$

بنابراین اگر تعداد نمونه‌های استخراج شده از متغیر تصادفی X بیستر از ۳ باشد خواهیم داشت:

$$\text{var}(\hat{\mu}_1) < \text{var}(\hat{\mu}_2)$$

و در نتیجه با توجه به معیار کارایی می‌توان نتیجه گرفت که برآورد کننده $\bar{X}$ بهتر از $\frac{1}{3} (X_1 + X_3 + X_5)$ در این مثال می‌باشد.

تعریف: برای دو برآوردگر $\hat{\theta}_1, \hat{\theta}_2$ نسبت $\frac{\text{var}(\hat{\theta}_1)}{\text{var}(\hat{\theta}_2)}$ را کارایی نسبی برآوردگر $\hat{\theta}_2$ نسبت به برآوردگر $\hat{\theta}_1$ می‌نامند.

بوضوح اگر مقدار کسر از یک بیشتر باشد می‌گوییم برآوردگر $\hat{\theta}_2$ از برآوردگر $\hat{\theta}_1$ کاراتر یا موثرتر است.

S-۷

مثال ۲۵: ۱۱: از یک جامعه نمونه‌ای به حجم ۲۰ با میانگین $\bar{X}_1 = 15$ و واریانس $\sigma_1^2 = 4$ استخراج کرده‌ایم. همچنین نمونه دیگری به حجم ۳ با

میانگین $\bar{X}_2 = 180$ و واریانس $\sigma_2^2 = 2$ استخراج کرده‌ایم. اگر $\bar{X}$ را به عنوان برآوردگر میانگین جامعه انتخاب کنیم کارایی نسبی $\bar{X}_1$ را محاسبه کنید؟

حل: ابتدا واریانس $\bar{X}_1$ و $\bar{X}_2$ را محاسبه می‌کنیم:

$$n_1 = 20, \quad \bar{X}_1 = 15, \quad \sigma_1^2 = 4$$

$$n_2 = 30, \quad \bar{X}_2 = 18, \quad \sigma_2^2 = 2$$

$$\text{var}(\bar{X}_1) = \frac{\sigma_1^2}{n_1} = \frac{4}{20} = \frac{1}{5}$$

$$\text{var}(\bar{X}_2) = \frac{\sigma_2^2}{n_2} = \frac{2}{30} = \frac{1}{15}$$

$$\bar{X}_2 \text{ نسبت به } \bar{X}_1 \text{ نسبی کارایی} = \frac{\text{var}(\bar{X}_2)}{\text{var}(\bar{X}_1)} = \frac{\frac{1}{15}}{\frac{1}{5}} = \frac{5}{15} = \frac{1}{3}$$

بنابراین استفاده از نمونه متوسط استخراج شده دوم برای برآورد میانگین جامعه کارایی بیشتری دارد و موثرتر می‌باشد. همچنین توجه کنید که در محاسبه کارایی از متوسط دو نمونه استخراج شده استفاده‌ای نمی‌شود.

۲۶ تا بحال نشان دادیم که یک برآوردگر به شرطی قابل قبول است که در درجه اول ناریب باشد و همچنین کمترین واریانس را در میان سایر برآوردگرها داشته باشد بنابراین تعریف زیر را داریم:

تعریف: اگر $\hat{\theta}$ یک برآورد کننده ناریب برای θ باشد. در این صورت $\hat{\theta}$ بهترین برآورد کننده ناریب خطی برای θ می‌باشد. هرگاه:

۱- $\hat{\theta}$ تابعی خطی از $X_1, X_2, \dots, X_n$ باشد.

۲- برای هر θ $E[\hat{\theta}] = \theta$

۳- از میان تمام برآورد کننده‌های $\hat{\theta}_i$ که در شرایط ۱ و ۲ صدق می‌کند داشته باشیم:

$$\text{var}(\hat{\theta}) \leq \text{var}(\hat{\theta}_i)$$

با توجه به استفاده از روش متوسط مربع خطا (MSE) تعریف بالا و شرایط آن کاملاً موجه می‌باشد. همچنین توجه کنید که اگر $\hat{\theta}$ در شرایط فوق صدق کند در این صورت $\hat{\theta}$ موثرترین برآورد کننده ناریب خطی نیز می‌باشد.
حال به مثال زیر توجه کنید:

۲۷ مثال ۱۲: فرض کنید X یک متغیر تصادفی با میانگین μ و واریانس σ^2 باشد. هرگاه $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از X باشد

نشان دهید توزیع متغیر تصادفی X هر چه باشد، برآورد کننده $\bar{X}$ بهترین برآورد کننده ناریب خطی برای μ می‌باشد؟

حل: ابتدا فرض می‌کنیم $\hat{\theta}$ یک برآورد کننده ناریب خطی برای μ باشد در این صورت $\hat{\theta}$ می‌بایستی به فرم زیر باشد:

$$\hat{\theta} = \sum_{i=1}^n (a_i x_i) + b$$

حال متغیرهای a_i و b را طوری بدست می‌آوریم که $\hat{\theta}$ بهترین برآورد کننده برای μ باشد. برای این منظور مطابق تعریف $\hat{\theta}$ می‌بایستی در شرایط تعریف بهترین برآورد کننده ناریب خطی صدق کند. یعنی: $E[\hat{\theta}] = \mu$ بنابراین:

$$\begin{aligned}
E[\hat{\theta}] &= E\left[\sum_{i=1}^n (a_i x_i) + b\right] = E\left[\sum_{i=1}^n a_i x_i\right] + b \\
&= E\sum_{i=1}^n (a_i E[X_i]) + b = \sum_{i=1}^n (a_i m_1) + b \\
&= m_1 \sum_{i=1}^n a_i + b = E[X] \sum_{i=1}^n a_i + b \Rightarrow E[\hat{\theta}] = \mu \sum_{i=1}^n a_i + b
\end{aligned}$$

حال چون بایستی داشته باشیم $E[\hat{\theta}] = v$ بنابراین:

$$\mu \sum_{i=1}^n a_i + b = \mu \Rightarrow \begin{cases} \sum_{i=1}^n a_i = 1 \\ b = 0 \end{cases}$$

۲۸ حال می‌بایستی واریانس برآورد کننده $\hat{\theta}$ را بدست آورده و سپس آنرا حداقل کنیم:

$$\begin{aligned}
\text{var}(\hat{\theta}) &= \text{var}\left(\sum_{i=1}^n a_i X_i + b\right) = \text{var}\left(\sum_{i=1}^n a_i X_i\right) \\
&= \sum \text{var}(a_i X_i) = \sum a_i^2 \text{var}(X_i) = \sum a_i^2 \sigma^2 = \sigma^2 \sum_{i=1}^n a_i^2
\end{aligned}$$

بنابراین برای حداقل نمودن واریانس $\sum a_i^2$ می‌بایستی حداقل شود به عبارتی باید مقادیر $a_1, a_2, \dots, a_n$ را طوری پیدا نمود که در دو شرط زیر صدق کنند:

$$\sum_{i=1}^n a_i = 1 \quad -1$$

$$\sum_{i=1}^n a_i^2 \quad -2 \text{ مینیمم شود.}$$

برای حداقل نمودن $\sum a_i^2$ ابتدا مقدار a_1 را از شرط اول محاسبه نموده و حاصل را در شرط دوم $(\sum a_i^2)$ قرار می‌دهیم. به صورت زیر:

$$\sum_{i=1}^n a_i = 1 \Rightarrow a_1 + \sum_{i=2}^n a_i = 1 \Rightarrow a_1 = 1 - \sum_{i=2}^n a_i$$

فرض می‌کنیم $Q = \sum_{i=1}^n a_i^2$ در این صورت باید Q مینیمم شود پس:

$$Q = \sum_{i=1}^n a_i^2 = a_1^2 + \sum_{i=2}^n a_i^2 = \left(1 - \sum_{i=2}^n a_i\right)^2 + \sum_{i=2}^n a_i^2$$

$$\Rightarrow Q = \left(1 - \sum_{i=2}^n a_i\right)^2 + \sum_{i=2}^n a_i^2$$

حال برای حداقل نمودن Q می‌بایستی مشتقات پاره‌ای Q را نسبت به a_i ها محاسبه نموده و برابر صفر قرار دهیم تا مقادیری بشکل

$a_1^*, a_2^*, \dots, a_n^*$ که Q را مینیمم می‌کنند بدست آید:

$$\frac{\partial Q}{\partial a_j} = 2(-1)\left(1 - \sum_{i=2}^n a_i\right) + 2a_j \quad j = 2, 3, \dots, n$$

۲۹ توجه کنید که مقدار a_1^* از روی سایر a_i^* ها ($i=2,3,\dots,n$) بشکل زیر بدست می‌آید:

$$a_1^* = 1 - \sum_{i=2}^n a_i^*$$

$$\frac{\partial Q}{\partial a_j} = 0 \Rightarrow \begin{cases} a_j^* = (1 - \sum_{i=2}^n a_i^*) = a_1^* \\ a_1^* = 1 - \sum_{i=2}^n a_i^* \end{cases}$$

رابطه ۱-۱۰

حال با حل دستگاه فوق مقادیر a_j^* ها بدست می‌آید، برای حل ابتدا کل معادلات را جداگانه نوشته و طرفین معادلات را جمع می‌کنیم. به این ترتیب داریم:

$$a_1^* = 1 - \sum_{i=2}^n a_i^*$$

$$a_2^* = 1 - \sum_{i=2}^n a_i^*$$

$$a_3^* = 1 - \sum_{i=2}^n a_i^*$$

⋮

$$a_n^* = 1 - \sum_{i=2}^n a_i^*$$

$$\text{حاصل جمع معادلات} = \sum_{i=1}^n a_i^* = n - n \sum_{i=2}^n a_i^* = n \underbrace{\left(1 - \sum_{i=2}^n a_i^*\right)}_{a_1^*} = n a_1^*$$

بنابراین بدست می‌آید:

$$\begin{cases} \sum_{i=1}^n a_i^* = n a_1^* \\ \sum_{i=1}^n a_i^* = 1 \end{cases} \Rightarrow n a_1^* = 1 \Rightarrow a_1^* = \frac{1}{n}$$

و با توجه به رابطه ۱-۱۰ بدست می‌آید:

$$(j=1,2,\dots,n) \quad a_j^* = \frac{1}{n}$$

بنابراین $\sum_{i=1}^n a_i^2$ و با شرط $\sum_{i=1}^n a_i = 1$ به ازای $a_j^* = \frac{1}{n}$ حداقل می‌شود. بنابراین برآوردگر $\hat{\theta}$ برابر می‌شود با: $\hat{\theta} = \sum_{i=1}^n \frac{1}{n} X_i = \frac{1}{n} \sum_{i=1}^n X_i = \bar{X}$

در نتیجه $\bar{X}$ بهترین برآوردگر ناریب خطی برای μ می‌باشد و این مساله مستقل از توزیع متغیر تصادفی X می‌باشد. به همین دلیل این مطلب از اهمیت فراوانی برخوردار است از طرفی می‌توان نشان داد که اگر برآوردگر $\hat{\theta}$ تابعی خطی از X_i ها هم نباشد باز هم $\bar{X}$ دارای کوچکترین واریانس می‌باشد. حال می‌توان گفت که $\bar{X}$ برای پارامتر μ در تابع نرمال و برای پارامتر p در توزیع برنولی و برای پارامتر λ در توزیع پواسون بهترین برآوردگر ناریب می‌باشد.

۳۰. ۳ برآورد کننده سازگار

یکی دیگر از خواصی که برای برآورد کننده‌ها می‌توان بدست آورد خاصیتی است که به سازگاری معروف است. اگر داشته باشیم:

$$\lim_{n \rightarrow \infty} \text{MSE}(\hat{\theta}) = \lim_{n \rightarrow \infty} E[(\hat{\theta} - \theta)^2] = 0$$

یعنی به ازای تعداد زیادی نمونه از متغیر تصادفی X مقدار متوسط مربع خطا برابر صفر شود، از طرفی طبق نامساوی مارکف (رجوع کنید به قضیه چبی شف) داریم:

$$p(g(x) \geq k) \leq \frac{E[g(x)]}{k} \Rightarrow 1 - p(g(x) \leq k) \leq \frac{E[g(x)]}{k}$$

$$\Rightarrow p(g(x) \geq k) \geq 1 - \frac{E[g(x)]}{k}$$

حال اگر $k = \varepsilon^2$ ($\varepsilon > 0$)، $g(x) = (\hat{\theta} - \theta)^2$ باشد در این صورت نامساوی بصورت زیر خواهد بود:

$$p((\hat{\theta} - \theta)^2 \leq \varepsilon^2) \geq 1 - \frac{E[(\hat{\theta} - \theta)^2]}{\varepsilon^2}$$

$$\Rightarrow p(|\hat{\theta} - \theta| < \varepsilon) \geq 1 - \frac{E[(\hat{\theta} - \theta)^2]}{\varepsilon^2}$$

حال اگر شرط $\lim_{n \rightarrow \infty} \text{MSE} = 0$ برقرار باشد می‌بایستی $E[(\hat{\theta} - \theta)^2]$ برابر صفر باشد در این صورت نامساوی بصورت زیر خلاصه می‌شود:

$$p(|\hat{\theta} - \theta| < \varepsilon) \geq 1 - \frac{0}{\varepsilon^2} \Rightarrow p(|\hat{\theta} - \theta| < \varepsilon) \geq 1$$

و از آنجا که احتمال حداکثر برابر ۱ واحد می‌باشد بنابراین:

$$p(|\hat{\theta} - \theta| < \varepsilon) = 1 \quad \forall \varepsilon > 0$$

$$p(|\hat{\theta} - \theta| \geq \varepsilon) = 0 \quad \forall \varepsilon > 0$$

یا به بیان دیگر:

در صورتی که احتمالات فوق برای یک برآورد کننده $\hat{\theta}$ برقرار باشند می‌گوییم برآورد کننده دارای خاصیت سازگاری می‌باشد.

بوضوح خاصیت سازگاری زمانی که تعداد نمونه‌ها به اندازه کافی بزرگ باشد نشان دهنده برتری یک برآورد کننده نسبت به دیگری می‌باشد. اما در

حالتی که تعداد نمونه‌ها کم باشد، سازگار بودن یک برآورد کننده معیاری برای بهتر بودن آن نمی‌باشد.

قبل از ارایه یک مثال به قضیه زیر نیز توجه کنید:

۳۱ قضیه: اگر داشته باشیم: $\lim_{n \rightarrow \infty} E[\hat{\theta}] = \theta$ و $\lim_{n \rightarrow \infty} \text{var}[\hat{\theta}] = 0$ آنگاه $\hat{\theta}$ یک برآورد کننده سازگار برای θ می‌باشد. این قضیه از آنجا

نتیجه می‌شود که در تعریف سازگاری داشتیم:

$$\lim_{n \rightarrow \infty} \text{MSE} = 0 \Rightarrow \lim_{n \rightarrow \infty} E[(\hat{\theta} - \theta)^2] = 0$$

$$\Rightarrow \lim_{n \rightarrow \infty} (\text{var}(\hat{\theta}) + (E[\hat{\theta}] - \theta)^2) = 0 \Rightarrow \begin{cases} \lim_{n \rightarrow \infty} \text{var}(\hat{\theta}) = 0 \\ \lim_{n \rightarrow \infty} E[\hat{\theta}] = \theta \end{cases}$$

۳۲ مثال ۱۳: اگر X یک متغیر تصادفی با میانگین μ و واریانس σ^2 باشد نشان دهید توزیع X هر چه باشد $\bar{X}$ یک برآورد کننده سازگار از روی نمونه‌های $X_1, X_2, \dots, X_n$ برای μ می‌باشد؟

حل: می‌دانیم:

$$E[\bar{X}] = \mu$$

$$\text{var}(\bar{X}) = \frac{\sigma^2}{n}$$

به این ترتیب:

$$\lim_{n \rightarrow \infty} E[X] = \lim_{n \rightarrow \infty} \mu = \mu$$

$$\lim_{n \rightarrow \infty} \text{var}(\bar{X}) = \lim_{n \rightarrow \infty} \frac{\sigma^2}{n} = 0$$

در نتیجه $\bar{X}$ یک برآورد کننده سازگار برای μ می‌باشد.

۳۳ مثال: اگر نمونه‌های $X_1, X_2, \dots, X_n$ را از یک جامعه نرمال با میانگین μ و واریانس σ^2 استخراج کنیم در این صورت نشان دهید که S^2

یک برآورد کننده سازگار برای σ^2 می‌باشد؟

حل: قبلاً نشان دادیم که $\frac{(n-1)S^2}{\sigma^2}$ یک متغیر تصادفی خی دو با $n-1$ درجه آزادی است (χ_{n-1}^2) و بنابراین میانگین آن $n-1$ و واریانس آن $2(n-1)$ می‌باشد:

$$E\left[\frac{(n-1)S^2}{\sigma^2}\right] = n-1$$

$$\text{var}\left(\frac{(n-1)S^2}{\sigma^2}\right) = 2(n-1)$$

حال با ساده نمودن روابط فوق داریم:

$$\frac{n-1}{\sigma^2} E[S^2] = n-1 \Rightarrow E[S^2] = \frac{n-1}{n-1} \sigma^2 \Rightarrow E[S^2] = \sigma^2$$

$$\text{var}\left(\frac{(n-1)S^2}{\sigma^2}\right) = 2(n-1) \Rightarrow \frac{(n-1)^2}{\sigma^4} \text{var}(S^2) = 2(n-1)$$

$$\text{var}(S^2) = \frac{2\sigma^4}{n-1} \Rightarrow \lim_{n \rightarrow \infty} \text{var}(S^2) = \lim_{n \rightarrow \infty} \frac{2\sigma^4}{n-1} = 0$$

بنابراین S^2 یک برآورد کننده سازگار برای σ^2 می‌باشد.

S-۹

۳۴ . ۱۰ ۴ آماره کافی

اگر بتوان یک آماره یافت که تمام اطلاعات را در بازه یک پارامتر مجهول خلاصه کند آنگاه آنرا یک آماره کافی برای پارامتر مجهول می‌نامیم. آماره‌های کافی به صورت زیر تعریف می‌شود:

تعریف: اگر نمونه تصادفی $X_1, X_2, \dots, X_n$ را از یک متغیر تصادفی با پارامتر مجهول θ داشته باشیم در این صورت آماره $y=g(X_1, X_2, \dots, X_n)$ را یک آماره کافی برای θ می‌نامیم اگر و فقط اگر برای هر مقدار $Y=y$ توزیع مشروط هر آماره بفرم $h(X_1, X_2, \dots, X_n)$ و با شرط $Y=y$ به θ وابسته نباشد.

به این معنی که اگر مقدار آماره کافی معلوم باشد ($Y=y$) آنگاه به خود مقادیر نمونه نیازی نیست و مقادیر نمونه اطلاعات بیشتری در بازه پارامتر مجهول θ نسبت به آماره کافی بدست آمده، نمی‌دهند.

مثال ۳۵: ۱۴: فرض می‌کنیم X_1, X_2 یک نمونه از متغیر تصادفی برنولی با پارامتر مجهول p باشند در این صورت نشان دهید که $Y = X_1 + X_2$ یک آماره کافی برای p می‌باشد؟
حل: تابع احتمال متغیر تصادفی برنولی بفرم زیر است:

$$p_X(x) = p^x (1-p)^{1-x} \quad x = 0, 1$$

سایر مقادیر

حال تعریف می‌کنیم:

$Y = X_1 + X_2 = g(X_1, X_2)$ باید نشان دهیم اگر $V = h(X_1, X_2)$ یک آماری دیگر باشد در این صورت $P_{V,Y}(V|Y=y)$ به θ وابسته نمی‌باشد. برای این منظور ابتدا لازم است توزیع توام V و Y را محاسبه کنیم که عبارتست از:

$$P_{V,Y}(V=h(0,0), Y=0) = (1-p)(1-p) = (1-p)^2$$

$$P_{V,Y}(V=h(0,1), Y=1) = p(1-p)$$

$$P_{V,Y}(V=h(1,0), Y=0) = (1-p)p$$

$$P_{V,Y}(V=h(1,1), Y=2) = p.p = p^2$$

حال توجه کنید که توزیع $Y = X_1 + X_2$ یک توزیع دو جمله‌ای با پارامتر $n=2$ و p می‌باشد بنابراین:

$$P_Y(y) = \binom{2}{y} p^y (1-p)^{2-y}$$

مقدار توزیع مشروط برای هر یک از مقادیر $Y=y$ عبارتست از:

$$P_{V,Y}(V=h(0,0) | Y=0) = \frac{(1-p)^2}{\binom{2}{0} p^0 (1-p)^{2-0}} = 1$$

$$P_{V,Y}(V=h(0,1) | Y=1) = \frac{p(1-p)}{\binom{2}{1} p(1-p)} = \frac{1}{2}$$

$$P_{V,Y}(V=h(1,1) | Y=2) = \frac{p^2}{\binom{2}{2} p^2 (1-p)^0} = 1$$

بنابراین اگر V هر آماره دلخواهی باشد، تابع احتمال مشروط آن به p وابسته نمی‌باشد در نتیجه $Y = X_1 + X_2$ یک آماره کافی برای پارامتر مجهول p می‌باشد.

۳۶ محاسبه کافی بودن یک آماره با توجه به تعریف کار بسیار مشکلی است، اما با استفاده از روش تجزیه فیشر-نیمن یک راه نسبتاً ساده برای اثبات کافی بودن یک آماره وجود دارد:

قضیه: $\hat{\theta}$ را یک برآورگر کافی θ می‌نامیم اگر و فقط اگر تابع چگالی احتمال توام نونه‌ها را بتوان بفرم زیر تجزیه کرد:

$$f(x_1, x_2, \dots, x_n, \theta) = \prod_{i=1}^n f_X(x_i) = g(\hat{\theta}, \theta) \cdot h(x_1, x_2, \dots, x_n)$$

مثال ۱۵: ۱۵: اگر یک نمونه n تایی از یک متغیر تصادفی برنولی داشته باشیم نشان دهید $Y = \sum_{i=1}^n X_i$ یک آماره کافی برای p می‌باشد؟

حل:

$$\prod_{i=1}^n P_X(x_i) = \prod_{i=1}^n [p^{x_i} (1-p)^{1-x_i}] = p^{\sum x_i} (1-p)^{n-\sum x_i}$$

$$\begin{cases} g(\hat{\theta}, \theta) = p^{\sum x_i} (1-p)^{1-x_i} \\ h(x_1, x_2, \dots, x_n) = 1 \end{cases}$$

حال داریم:

بنابراین $Y = \sum_{i=1}^n X_i$ یک آماره کافی برای p می‌باشد.

۳۷ توجه کنید هرگاه $\hat{\theta}$ یک برآورد کننده نااریب برای θ بر حسب یک آماره کافی باشد آنگاه می‌توان نشان داد که $\hat{\theta}$ دارای کمترین واریانس در میان سایر برآورد کننده‌های نااریب برای θ می‌باشد. به همین دلیل است که برآورد کننده‌ها را بر پایه آماره‌های کافی بدست می‌آورند.

قبلاً نشان دادیم که $S^2 = \left(\frac{n-1}{n}\right)$ یک برآورد کننده نااریب برای σ^2 متغیر تصادفی نرمال می‌باشد، حال با استفاده از برآورد کننده کافی یک برآورد کننده نااریب و کافی برای متغیر تصادفی نرمال بدست می‌آوریم.

برای متغیر تصادفی نرمال با میانگین μ و واریانس σ^2 داریم:

$$\prod_{i=1}^n f_X(x_i) = \prod_{i=1}^n \left(\frac{1}{\sigma \sqrt{2\pi}} \right) e^{-\frac{(x_i - \mu)^2}{2\sigma^2}} = \left(\frac{1}{\sigma} \right)^n \left(\frac{1}{2\pi} \right)^{\frac{n}{2}} e^{-\frac{\sum (x_i - \mu)^2}{2\sigma^2}}$$

$$= g(\hat{\theta}, \theta) \cdot h(x_1, x_2, \dots, x_n)$$

$$h(x_1, x_2, \dots, x_n) = \left(\frac{1}{2\pi} \right)^{\frac{n}{2}}$$

که در آن:

از طرفی به جای μ نیز می‌توان از برآورد کافی آن $\bar{X}$ استفاده نمود. در نتیجه $\sum (x_i - \bar{X})^2$ یک آماره کافی برای σ^2 می‌باشد. از طرفی قبلاً نشان دادیم که:

$$E\left[\frac{1}{n} \sum (x_i - \bar{X})^2 \right] = \frac{n-1}{n} \sigma^2$$

$$\Rightarrow E\left[\frac{n-1}{n} \underbrace{\left(\frac{1}{n-1} \sum (x_i - \bar{X})^2 \right)}_{S^2} \right] = E\left[\frac{n-1}{n} S^2 \right] = \frac{n-1}{n} \sigma^2$$

$$\Rightarrow \frac{n-1}{n} E[S^2] = \frac{n-1}{n} \sigma^2 \Rightarrow E[S^2] = \sigma^2$$

بنابراین خود S^2 یک برآورد کننده نااریب و کافی برای σ^2 می‌باشد. و دارای کوچکترین واریانس در میان چنین برآورد کننده‌هایی می‌باشد.

۵.۱۰.۳۸ خاصیت پایایی برآورد کننده‌ها

یکی از نتایج کلی که از برآورد کننده‌های حداکثر احتمال بدست می‌آید خاصیت پایایی برآورد کننده حداکثر احتمال است. اگر $\hat{\theta}$ برآورد کننده حداکثر احتمال θ باشد و بخواهیم تابعی از θ مثل $\gamma = h(\theta)$ را برآورد کنیم کافیست به جای θ مقدار برآورد شده آنرا که $\hat{\theta}$ می‌باشد جایگزین کنیم در این صورت برآورد γ عبارتست از:

$$\hat{\gamma} = h(\hat{\theta})$$

این مطلب نشان می‌دهد که برای محاسبه برآورد تابعی از پارامتر مجهول نیازی به محاسبه مجدد برآورد کننده نمی‌باشد بلکه کافیست یکبار برآورد برای پارامتر θ بدست آید. توجه کنید که این مطلب فقط برای برآورد کننده‌هایی که به روش حداکثر احتمال بدست می‌آیند صادق می‌باشد.

مثال ۱۶: اگر X یک متغیر تصادفی هندسی نوع اول باشد مطلوبست:

الف) برآورد پارامتر p به روش حداکثر احتمال.

ب) برآورد تابع $h(p) = p(X \leq 2)$ با استفاده از خاصیت پایایی.

حل: تابع احتمال توزیع هندسی نوع اول عبارتست از:

$$P_X(x) = p(1-p)^x \quad x=0,1,2,\dots$$

= 0 سایر مقادیر

$$L(p) = \prod_{i=1}^n p(1-p)^{x_i} = p^n (1-p)^{\sum x_i}$$

$$H(p) = \ln(L(p)) = n \ln(p) + \sum x_i \ln(1-p)$$

$$\frac{dH}{dp} = \frac{n}{p} - \frac{\sum x_i}{1-p} \Rightarrow \frac{1-\hat{p}}{\hat{p}} = \frac{1}{n} \sum x_i = \bar{X} \Rightarrow \frac{1}{\hat{p}} - 1 = \bar{X} \Rightarrow \hat{p} = \frac{1}{1+\bar{X}}$$

بنابراین برآورد p به روش حداکثر احتمال برابر $\frac{1}{1+\bar{X}}$ می باشد.

ب)

$$p(X \leq 2) = p(X=0) + p(X=1) + p(X=2)$$

$$= p + p(1-p) + p(1-p)^2$$

$$\gamma = h(p) = p(1 + (1-p) + (1-p)^2)$$

حال داریم:

بنابراین خاصیت پایایی $\hat{\gamma} = h(\hat{p})$ بنابراین:

$$\hat{\gamma} = h(\hat{p}) = \frac{1}{1+\bar{X}} \left(1 + \left(1 - \frac{1}{1+\bar{X}} \right) + \left(1 - \frac{1}{1+\bar{X}} \right)^2 \right)$$

S-۱۰

۱۰۴۰. ۶ برآورد فاصله‌ای

در برآورد نقطه‌ای با توجه به نمونه‌های بدست آمده نشان دادیم برای پارامترهای مجهول جامعه می‌توان عددی بدست آورد که با خطای کمی بیانگر مقدار واقعی پارامتر مجهول باشد. اما واقعاً نمی‌دانیم با چه مقدار خطایی پارامتر مجهول را برآورد کرده‌ایم برآورد فاصله‌ای روشی است که در آن برای پارامتر مجهول یک بازه در نظر می‌گیریم و نشان می‌دهیم که با احتمالی معین پارامتر مجهول در بازه مورد نظر قرار دارد.

فرض کنید نمونه‌های $X_1, X_2, \dots, X_n$ را از یک متغیر تصادفی X بدست آورده باشیم در این صورت فاصله (L_1, L_2) را طوری می‌سازیم که با احتمال بسیار زیاد مثلاً ۹۵٪ شامل پارامتر مجهول جامعه (مثلاً μ) باشد. به فاصله (L_1, L_2) یک فاصله اطمینان می‌گویند زیرا مطمئن هستیم که به احتمال زیاد (مثلاً ۹۵٪) این فاصله شامل پارامتر مجهول می‌باشد. مقادیر L_1, L_2 هر کدام آماره‌ای هستند که بر اساس مقادیر معلوم نمونه‌ها بدست می‌آیند. حال به تعریف دقیق فاصله اطمینان توجه کنید:

هرگاه مقادیر $X_1, X_2, \dots, X_n$ تعداد n نمونه تصادفی بدست آمده از یک متغیر تصادفی X با پارامتر مجهول θ باشند. در این صورت می‌گوییم دو آماره L_1, L_2 یک $(1-\alpha)$ ۱۰۰ درصد فاصله اطمینان برای θ تشکیل می‌دهند. هرگاه:

$$p(L_1 \leq \theta \leq L_2) \geq 1-\alpha$$

توجه کنید که $1-\alpha$ مقدار احتمال می‌باشد و بنابراین در بازه ۰ تا ۱ قرار دارد، در نتیجه برای بیان آن بصورت درصد آنرا درصد ضرب کرده‌ایم و به صورت $(1-\alpha)$ ۱۰۰٪ می‌نویسیم.

در برآورد فاصله‌ای از آماره‌های بدست آمده از فصل قبل استفاده می‌کنیم. اینکه کدام پارامتر جامعه مجهول باشد، تعیین می‌کند که از کدام آماره در ساختن فاصله اطمینان استفاده کنیم. مثال بعد روش ساختن فاصله اطمینان برای جامعه نرمال را نشان می‌دهد.

۴۱ مثال ۱۷: در شهر A بارش باران یک متغیر تصادفی نرمال با میانگین μ سانتی متر و واریانس ۴ می باشد اگر میزان بارش را در طول ۱۶ روز اندازه بگیریم و میانگین ۱۰ سانتی متر را بدست بیاوریم، در این صورت با احتمال ۰.۹۵ میانگین بارش باران در شهر A در چه فاصله اطمینانی قرار می گیرد؟
حل: معلومات مساله عبارتند از:

$$n = 16$$

$$\bar{X} = 10$$

$$\mu = \text{مجهول}$$

$$\sigma^2 = 4 \rightarrow \sigma = 2$$

ابتدا توجه کنید که با توجه به جدول متغیر تصادفی نرمال استاندارد رابطه زیر برقرار می باشد:

$$p(-1/96 \leq Z \leq 1/96) = 0.95 \quad (\text{رابطه } 10-2)$$

زیرا داریم:

$$N_Z(1/96) = p(Z \leq 1/96) = 0.975$$

$$N_Z(-1/96) = 1 - N_Z(1/96) = 0.025$$

$$\Rightarrow p(-1/96 \leq Z \leq 1/96) = p(Z \leq 1/96) - p(Z \leq -1/96) = 0.975 - 0.025 = 0.95$$

۴۲ در فصل نمونه گیری نشان دادیم که اگر واریانس جامعه معلوم باشد آماره $Z = \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}}$ یک متغیر تصادفی نرمال استاندارد می باشد و از طرفی با توجه به رابطه ۱۰ - ۲ چون Z نیز یک متغیر تصادفی نرمال استاندارد می باشد می توان نوشت:

$$P(-1/96 \leq Z \leq 1/96) = P(-1/96 \leq \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} \leq 1/96)$$

$$\text{با ساده نمودن نامساوی} \quad -1/96 \leq \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} \leq 1/96 \quad \text{بدست می آوریم:}$$

$$\bar{X} - 1/96 \left(\frac{\sigma}{\sqrt{n}} \right) \leq \mu \leq \bar{X} + 1/96 \left(\frac{\sigma}{\sqrt{n}} \right)$$

بنابراین رابطه ۱۰ - ۲ بصورت زیر بدست می آید:

$$p\left(\bar{X} - 1/96 \left(\frac{\sigma}{\sqrt{n}} \right) \leq \mu \leq \bar{X} + 1/96 \left(\frac{\sigma}{\sqrt{n}} \right)\right) = 0.95$$

همانطور که ملاحظه می کنید برای μ یک فاصله اطمینان ۹۵ درصدی به شکل (L_1, L_2) بدست آوردیم که در آن L_1, L_2 عبارتند از:

$$L_1 = \bar{X} - 1/96 \left(\frac{\sigma}{\sqrt{n}} \right)$$

$$L_2 = \bar{X} + 1/96 \left(\frac{\sigma}{\sqrt{n}} \right)$$

حال با جایگذاری $\bar{X}$ و σ و n در رابطه بدست می آوریم:

$$L_1 = 10 - 1/96 \left(\frac{2}{\sqrt{16}} \right) = 9.02$$

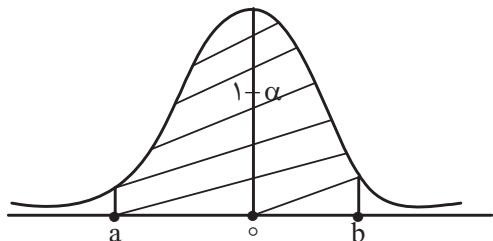
$$L_2 = 10 + 1/96 \left(\frac{2}{\sqrt{16}} \right) = 10.98$$

$$p(9.02 \leq \mu \leq 10.98) = 0.95$$

در نتیجه:

یعنی ۹۵ درصد اطمینان داریم که میانگین جامعه در بازه $(\frac{10}{98}, \frac{9}{102})$ قرار دارد حال برای اینکه بتوانیم یک فاصله اطمینان $(1-\alpha) \cdot 100$ درصدی در حالت کلی برای جامعه نرمال با میانگین مجهول و واریانس معلوم بسازیم ابتدا مفهوم $Z_{1-\frac{\alpha}{2}}$ را معرفی می‌کنیم.

۴۳ نمودار زیر منحنی نمایش توزیع نرمال استاندارد می‌باشد ملاحظه می‌کنید که مساحت مشخص شده در شکل برابر $1-\alpha$ می‌باشد:



برای پیدا کردن فاصله اطمینان می‌بایستی اعداد a و b را طوری پیدا کرد که:

$$p(a \leq Z \leq b) = p(a \leq \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} \leq b) = 1 - \alpha$$

با فرض اینکه فاصله متقارن باشد مساحت زیر منحنی بعد از نقطه b و قبل از نقطه a برابر $\frac{\alpha}{2}$ می‌باشد. در نتیجه مساحت کل زیر منحنی قبل از

نقطه b برابر است با $1 - \frac{\alpha}{2}$ که آنرا با $Z_{1-\frac{\alpha}{2}}$ نمایش می‌دهیم. به این ترتیب $Z_{1-\frac{\alpha}{2}}$ نقطه‌ای روی محور افقی است که مساحت

زیر منحنی قبل از آن برابر $1 - \frac{\alpha}{2}$ می‌باشد. این مطلب با توجه به تعریف تابع توزیع نرمال استاندارد بصورت زیر نیز قابل بیان است:

$$p(Z \leq Z_{1-\frac{\alpha}{2}}) = 1 - \frac{\alpha}{2}$$

حال می‌توان به جای a و b در $p(a \leq Z \leq b)$ مقادیر $Z_{1-\frac{\alpha}{2}}$ و $-Z_{1-\frac{\alpha}{2}}$ را جایگزین نمود و با ساده نمودن نامساوی خواهیم داشت:

$$p(-Z_{1-\frac{\alpha}{2}} \leq Z \leq Z_{1-\frac{\alpha}{2}}) = p(-Z_{1-\frac{\alpha}{2}} \leq \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} \leq Z_{1-\frac{\alpha}{2}})$$

$$p = (-Z_{1-\frac{\alpha}{2}} (\frac{\sigma}{\sqrt{n}}) \leq \bar{X} - \mu \leq Z_{1-\frac{\alpha}{2}} (\frac{\sigma}{\sqrt{n}}))$$

$$p = (\bar{X} - Z_{1-\frac{\alpha}{2}} (\frac{\sigma}{\sqrt{n}}) \leq \mu \leq \bar{X} + Z_{1-\frac{\alpha}{2}} (\frac{\sigma}{\sqrt{n}})) = 1 - \alpha$$

بنابراین با توجه به آخرین رابطه یک فاصله $(1-\alpha) \cdot 100$ درصد برای میانگین (μ) جامعه نرمال با معلوم بودن σ^2 بدست آوردیم.

۴۴ مثال ۱۸: میزان فروش ماهانه یک فروشگاه یک متغیر تصادفی نرمال با میانگین مجهول μ و انحراف معیار ۵۰ هزار تومان در ماه می‌باشد.

نتایج حاصل از ۲۵ ماه فروش نشان می‌دهند که میانگین فروش در طول این ۲۵ ماه عبارتست از ۵۵۰ هزار تومان. اگر ۹۰ درصد اطمینان داشته

باشیم میانگین (μ) در فاصله (a, b) می‌باشد مقادیر a و b را مشخص کنید؟

حل: معلومات مثال عبارتند از:

$$n = 25$$

$$\bar{X} = 55.0 \quad \mu = \text{مجهول} \quad \sigma = 5.0$$

می‌دانیم فاصله اطمینان از رابطه زیر بدست می‌آید:

$$p\left(\bar{X} - Z_{1-\frac{\alpha}{2}} \left(\frac{\sigma}{\sqrt{n}}\right) \leq \mu \leq \bar{X} + Z_{1-\frac{\alpha}{2}} \left(\frac{\sigma}{\sqrt{n}}\right)\right) = 1-\alpha$$

ابتدا می‌بایستی مقدار α را بدست بیاوریم و سپس از روی آن $Z_{1-\frac{\alpha}{2}}$ بدست می‌آید.

$$100\%(1-\alpha) = 90\% \Rightarrow 1-\alpha = 0.9 \Rightarrow \alpha = 0.1$$

$$Z_{1-\frac{\alpha}{2}}^{\alpha=0.1} = Z_{1-\frac{0.1}{2}} = Z_{0.95}$$

۴۵ با توجه به جدول توزیع نرمال استاندارد $Z_{0.95}$ عبارتست از:

$$N_Z(Z_{0.95}) = 0.95 \Rightarrow Z_{0.95} = 1.645$$

بنابراین بدست می‌آوریم:

$$p\left(55.0 - 1.645 \left(\frac{5.0}{\sqrt{25}}\right) \leq \mu \leq 55.0 + 1.645 \left(\frac{5.0}{\sqrt{25}}\right)\right) = 0.9$$

$$\Rightarrow p(53.3/55 \leq \mu \leq 56.6/45) = 0.9$$

بنابراین با احتمال ۰/۹ متوسط درآمد فروشگاه در بازه (۵۳۳/۵۵ , ۵۶۶/۴۵) قرار دارد. توجه کنید که طول بازه فاصله اطمینان در حالت کلی برابر است با:

$$L_2 - L_1 = \left(\bar{X} + Z_{1-\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}\right) - \left(\bar{X} - Z_{1-\frac{\alpha}{2}} \left(\frac{\sigma}{\sqrt{n}}\right)\right) \Rightarrow L_2 - L_1 = 2 Z_{1-\frac{\alpha}{2}} \left(\frac{\sigma}{\sqrt{n}}\right)$$

ملاحظه می‌کنید که طول بازه مستقل از $\bar{X}$ می‌باشد و تنها به تعداد نمونه‌ها وابسته می‌باشد. یعنی می‌توان تعداد نمونه‌ها را طوری انتخاب نمود که هر فاصله اطمینان دلخواهی برای μ بدست آید.

به عنوان مثال در مثال قبل با ۲۵ نمونه طول بازه عبارتست از:

$$L_2 - L_1 = 2(1.645) \left(\frac{5.0}{\sqrt{100}}\right) = 1.645$$

پیدا است که با ۱۰۰ نمونه طول بازه نصف شده است.

۱۰.۶.۱۰ فاصله اطمینان برای جامعه نرمال با میانگین مجهول μ و واریانس مجهول σ^2 . حال فاصله اطمینان را برای حالتی بدست می‌آوریم که

علاوه بر میانگین جامعه، واریانس نیز مجهول باشد. از فصل نمونه‌گیری می‌دانیم متغیر تصادفی $T = \frac{\bar{X} - \mu}{\frac{S}{\sqrt{n}}}$ یک متغیر تصادفی t با $n-1$ درجه

آزادی می‌باشد بنابراین مجدداً مشابه روشی که قبلاً برای بدست آوردن فاصله اطمینان ارایه کردیم داریم:

$$p\left(-t_{1-\frac{\alpha}{2}} \leq t_{(n-1)} \leq t_{1-\frac{\alpha}{2}}\right) = p\left(-t_{1-\frac{\alpha}{2}} \leq \frac{\bar{X} - \mu}{\frac{S}{\sqrt{n}}} \leq t_{1-\frac{\alpha}{2}}\right)$$

$$= p\left(\bar{X} - t_{1-\frac{\alpha}{2}} \left(\frac{S}{\sqrt{n}}\right) \leq \mu \leq \bar{X} + t_{1-\frac{\alpha}{2}} \left(\frac{S}{\sqrt{n}}\right)\right) = 1-\alpha$$

که در آن $t_{1-\frac{\alpha}{2}}$ عبارتست از: $t_{1-\frac{\alpha}{2}}(n-1) = 1 - \frac{\alpha}{2}$.

بنابراین فاصله اطمینان در این حالت عبارتست از:

$$L_1 = \bar{X} - t_{1-\frac{\alpha}{2}} \left(\frac{S}{\sqrt{n}} \right)$$

$$L_2 = \bar{X} + t_{1-\frac{\alpha}{2}} \left(\frac{S}{\sqrt{n}} \right)$$

۴۷ مثال ۱۹: میزان مصرف بنزین در هر ماه یک متغیر تصادفی نرمال با میانگین مجهول μ و واریانس σ^2 می‌باشد اگر اطلاعات زیر از یک نمونه در طول ۶ ماه بدست آمده باشد مطلوبست یک فاصله اطمینان ۹۰ درصدی برای متوسط میزان مصرف بنزین؟

$$\sum_{i=1}^6 x_i = 246 \quad \sum_{i=1}^6 x_i^2 = 11200$$

حل:

$$\bar{X} = \frac{1}{n} \sum_{i=1}^n x_i = \frac{1}{6} (246) = 41$$

$$S^2 = \frac{1}{n-1} \left(\sum_{i=1}^n (x_i - \bar{X})^2 \right) = \frac{n}{n-1} \left(\frac{1}{n} \sum_{i=1}^n (x_i - \bar{X})^2 \right)$$

$$= \frac{n}{n-1} \left(\frac{1}{n} \sum_{i=1}^n x_i^2 - \bar{X}^2 \right) = \frac{1}{n-1} \left(\sum_{i=1}^n x_i^2 - n \bar{X}^2 \right)$$

$$\Rightarrow S^2 = \frac{1}{5} (11200 - 6(41)^2) = \frac{1}{5} (1114) = 222.8$$

$$\Rightarrow S = 14.92$$

حال برای یک فاصله اطمینان ۹۰ درصدی داریم:

$$1 - \alpha = 0.9 \Rightarrow \alpha = 0.1$$

$$t_{1-\frac{\alpha}{2}} = t_{0.95}$$

۴۸ با توجه به جدول t با ۵ درجه آزادی مقدار $t_{0.95}$ برابر ۲/۰۱۵ بدست می‌آید.

در نتیجه:

$$L_1 = \bar{X} - t_{0.95} \left(\frac{S}{\sqrt{n}} \right) = 41 - 2.015 \left(\frac{14.92}{\sqrt{6}} \right) = 28.72$$

$$L_2 = \bar{X} + t_{0.95} \left(\frac{S}{\sqrt{n}} \right) = 41 + 2.015 \left(\frac{14.92}{\sqrt{6}} \right) = 53.27$$

بنابراین با ۹۰ درصد اطمینان می‌توانیم بگوییم که متوسط مصرف بنزین ماهانه مابین ۲۸/۷۲ و ۵۳/۲۷ میلیون لیتر خواهد بود.

برای حالتی که σ^2 مجهول باشد طول بازه فاصله اطمینان عبارتست از:

$$L_2 - L_1 = 2 t_{1-\frac{\alpha}{2}} \left(\frac{S}{\sqrt{n}} \right)$$

اما توجه می‌کنید که در این حالت طول بازه به متغیر S نیز وابسته می‌باشد و در نتیجه به هیچ عنوان نمی‌توان اندازه نمونه را طوری انتخاب نمود که دقیقاً یک بازه با طول معین بدست آید.

نکته: قبلاً نشان دادیم که اگر تعداد نمونه‌ها زیاد باشد ($n > 30$) $\frac{\bar{X}-\mu}{\frac{S}{\sqrt{n}}}$ و $\frac{\bar{X}-\mu}{\frac{\sigma}{\sqrt{n}}}$ هر دو به توزیع نرمال استاندارد میل می‌کنند، در نتیجه تمام مطالب گفته شده در حالتی که $n > 30$ باشد برای متغیر تصادفی X با هر توزیعی برقرار می‌باشد.

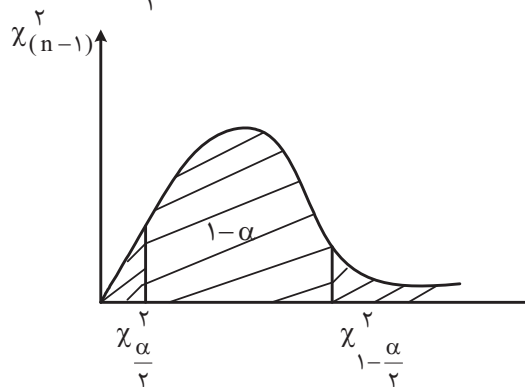
۴۹. ۱۰. ۶. ۲ فاصله اطمینان برای واریانس جامعه نرمال با میانگین مجهول μ و واریانس σ^2 .

روش بدست آوردن فاصله اطمینان برای واریانس یک جامعه نرمال با میانگین مجهول μ و واریانس σ^2 کاملاً مشابه روشی است که برای میانگین جامعه استفاده شد. در این حالت از آماره $\frac{(n-1) S^2}{\sigma^2}$ استفاده می‌کنیم. می‌دانیم که یک متغیر تصادفی χ^2 با $n-1$ درجه آزادی است. می‌دانیم:

$$p \left(\chi_{\frac{\alpha}{2}}^2 \leq \chi_{(n-1)}^2 \leq \chi_{1-\frac{\alpha}{2}}^2 \right) = 1-\alpha$$

$$\Rightarrow p \left(\chi_{\frac{\alpha}{2}}^2 \leq \frac{(n-1)S^2}{\sigma^2} \leq \chi_{1-\frac{\alpha}{2}}^2 \right) = 1-\alpha \quad (3-10)$$

توجه کنید که مقدار ابتدایی بازه (L_1) برابر $\chi_{\frac{\alpha}{2}}^2$ می‌باشد زیرا مطابق شکل منحنی توزیع $\chi_{(n-1)}^2$ مساحت زیر منحنی قبل از $\chi_{\frac{\alpha}{2}}^2$ برابر $\frac{\alpha}{2}$ می‌باشد.



$$p \left(\chi_{(n-1)}^2 \leq \chi_{\frac{\alpha}{2}}^2 \right) = \frac{\alpha}{2}$$

حال با ساده نمودن رابطه (۳-۱۰) داریم:

$$p \left(\frac{(n-1) S^2}{\chi_{1-\frac{\alpha}{2}}^2} \leq \sigma^2 \leq \frac{(n-1) S^2}{\chi_{\frac{\alpha}{2}}^2} \right) = 1-\alpha$$

بنابراین بازه زیر یک فاصله اطمینان $(1-\alpha)$ ۱۰۰ درصدی برای واریانس بدست می‌دهد:

$$L_1 = \frac{(n-1) S^2}{\chi_{1-\frac{\alpha}{2}}^2}, \quad L_2 = \frac{(n-1) S^2}{\chi_{\frac{\alpha}{2}}^2}$$

۵۰ مثال ۲۰: کارخانه‌ای ظروف پلاستیکی تولید می‌کند، ضایعات تولید شده توسط ماشین‌های تولیدی یک متغیر تصادفی نرمال با میانگین μ و

واریانس σ^2 می‌باشد، اگر پراکندگی ضایعات تولید شده در روز زیاد شود (واریانس تعداد ضایعات زیاد شود). می‌بایستی ماشین‌های تولیدی را تعمیر

نمود. برای تخمین واریانس یک نمونه ۲۵ تایی از تعداد ضایعات تولید شده در هر روز می‌گیریم و نتایج $\sum x_i = 170$ ، $\sum x_i^2 = 2750$ را بدست می‌آوریم.

در صورتی که یک فاصله اطمینان ۹۰ درصدی برای واریانس بدست بیاوریم و طول فاصله اطمینان بیشتر از ۵۰ شود می‌بایستی ماشینهای تولیدی را تعمیر کنیم، آیا نیازی به تعمیر می‌باشد؟
حل: معلومات مثال عبارتند از:

$$n = 25$$

$$\bar{X} = \frac{1}{25} (170) = 6.8$$

$$S^2 = \frac{1}{n-1} (\sum x_i^2 - n \bar{X}^2) = \frac{1}{24} (2750 - 25(6.8)^2) = 66/41$$

$$\Rightarrow S = 8/14$$

$$1-\alpha = 0.9 \rightarrow \alpha = 0.1 \rightarrow \chi_{\frac{\alpha}{2}}^2 (n-1) = \chi_{0.05}^2 (24) = 13/8$$

$$\chi_{1-\frac{\alpha}{2}}^2 (n-1) = \chi_{0.95}^2 (24) = 36/4$$

بنابراین:

$$L_1 = \frac{(n-1) S^2}{\chi_{1-\frac{\alpha}{2}}^2} = \frac{(24) (66/41)}{36/4} = 43/78$$

$$L_2 = \frac{(n-1) S^2}{\chi_{\frac{\alpha}{2}}^2} = \frac{24 (66/41)}{13/8} = 115/49$$

$$L_2 - L_1 = 115/49 - 43/78 = 71/7$$

با توجه به اینکه طول بازه فاصله اطمینان بیشتر از ۵۰ واحد شده است بنابراین می‌بایستی ماشینهای تولیدی تعمیر شوند.

۵۱. ۶. ۳ فاصله اطمینان برای متغیر تصادفی نمایی با پارامتر مجهول λ

می‌دانیم متغیر تصادفی نمایی عموماً مدلی برای طول عمر وسایل الکتریکی، تجهیزات و ... می‌باشد. بنابراین اگر بخواهیم اطمینان داشته باشیم که متوسط طول عمر یک متغیر تصادفی نمایی حداقل برابر تعداد ساعات معین است در این صورت می‌بایستی $\frac{1}{\lambda}$ از عدد معین بیشتر باشد بنابراین برای متغیر تصادفی نمایی تنها یک فاصله اطمینان یک طرفه می‌سازیم.

هرگاه $X_1, X_2, \dots, X_n$ یک نمونه تصادفی از یک متغیر تصادفی نمایی با پارامتر λ باشد آنگاه:

$$p(\lambda \leq \frac{\chi_{1-\alpha}^2}{2n\bar{X}}) = 1-\alpha$$

یک فاصله اطمینان $(1-\alpha)$ ۱۰۰ درصدی برای λ تشکیل می‌دهد که در آن $\chi_{1-\alpha}^2$ مقداری از توزیع χ^2 با $2n$ درجه آزادی است که مساحت زیر منحنی نمایش آن برابر $1-\alpha$ باشد.

۵۲ مثال ۲۱: طول عمر یک مدار الکتریکی یک متغیر تصادفی نمایی با پارامتر λ می باشد تعداد ۲۰ عدد از این مدارها را آزمایش می کنیم تا زمانی که از کار بیفتند. و ملاحظه می کنیم که مجموع طول آنها برابر ۲۳۶۰ ساعت می شود مطلوبست: تخمین متوسط طول عمر مدار مربوط با فاصله اطمینان ۹۰ درصدی:
حل: معلومات مساله عبارتند از:

$$n = 20$$

$$\sum_{i=1}^{20} x_i = 2360 \rightarrow \bar{X} = \frac{1}{n} \sum_{i=1}^n x_i = \frac{1}{20} (2360) = 118$$

$$1-\alpha = 0.9 \rightarrow \chi^2_{1-\alpha}(2n) = \chi^2_{0.9}(40) = 51.8$$

بنابراین یک فاصله اطمینان یک طرفه ۹۰ درصدی برای λ عبارتست از:

$$p\left(\lambda \leq \frac{\chi^2_{1-\alpha}}{2n\bar{X}}\right) = 1-\alpha \Rightarrow L_2 = \frac{51.8}{2(20)(118)} = 0.0109$$

توجه کنید که مقدار متوسط طول عمر برای توزیع نمایی $E[X] = \frac{1}{\lambda}$ می باشد بنابراین حداقل متوسط طول عمر مدار عبارتست از:

$$\frac{1}{0.0109} = 91.11$$

به عبارتی ۹۰ درصد اطمینان داریم که هر مدار حداقل ۹۱/۱۱ ساعت کار می کند.

S-۱۳

۵۳ ۱۰. ۶. ۴ فاصله اطمینان برای نسبت p با تعداد نمونه های زیاد در توزیع برنولی

نمونه های $X_1, X_2, \dots, X_n$ را از یک متغیر تصادفی برنولی با پارامتر مجهول p استخراج می کنیم اگر تعداد نمونه ها زیاد باشد در این صورت می دانیم متغیر تصادفی $\bar{X}$ تقریباً بصورت نرمال با میانگین p و واریانس $\frac{p(1-p)}{n}$ توزیع می شود. به این ترتیب می توان نشان داد که یک فاصله اطمینان $(1-\alpha) 100$ درصدی برای p عبارتست از:

$$p\left(\bar{X} - Z_{1-\frac{\alpha}{2}} \sqrt{\frac{\bar{X}(1-\bar{X})}{n}} \leq p \leq \bar{X} + Z_{1-\frac{\alpha}{2}} \sqrt{\frac{\bar{X}(1-\bar{X})}{n}}\right) = 1-\alpha$$

توجه کنید که در محاسبه مقدار $\bar{X}$ برابر است با نسبت تعداد موفقیتها به کل تعداد نمونه ها به عنوان مثال اگر از ۲۰ مرتبه تیراندازی به هدف ۱۵ بار هدف مورد اصابت قرار گیرد در این صورت $\bar{X} = \hat{p} = \frac{15}{20}$ خواهد بود.

۵۴ مثال ۲۲: احتمال اینکه یک بازیکن فوتبال یک ضربه پنالتی را گل کند یک متغیر تصادفی برنولی با پارامتر p می باشد. بازیکن مربوطه در تمرینات از تعداد ۲۰۰ ضربه پنالتی ۱۶۰ مرتبه را موفق به زدن گل شده است. اگر بخواهیم در مسابقات بازیکن مربوطه را مامور زدن ضربه پنالتی کنیم با چه احتمالی می توانیم ۹۵ درصد اطمینان داشته باشیم که وی گل می زند؟

حل: معلومات مساله عبارتند از:

$$n = 200$$

$$\bar{X} = \hat{p} = \frac{160}{200} = \frac{16}{20} = 0.8$$

$$1 - \alpha = 0.95 \rightarrow \alpha = 0.05 \rightarrow Z_{1-\frac{\alpha}{2}} = Z_{0.975} = 1.96$$

بنابراین یک فاصله ۹۵ درصدی برای p عبارتست از:

$$L_1 = \bar{X} - Z_{1-\frac{\alpha}{2}} \left(\sqrt{\frac{\bar{X}(1-\bar{X})}{n}} \right) = 0.8 - 1.96 \sqrt{\frac{0.8(0.2)}{200}} = 0.74$$

$$L_2 = \bar{X} + Z_{1-\frac{\alpha}{2}} \left(\sqrt{\frac{\bar{X}(1-\bar{X})}{n}} \right) = 0.8 + 1.96 \sqrt{\frac{0.8(0.2)}{200}} = 0.85$$

بنابراین با ۹۵ درصد اطمینان می‌توانیم بگوییم که احتمال گل شدن ضربه پنالتی توسط بازیکن مربوطه در بازه $(0.74, 0.85)$ قرار دارد.

۵۵. ۱۰. ۵. فاصله اطمینان برای تفاضل نسبت دو جامعه با تعداد نمونه‌های زیاد

اگر دو نمونه با حجم‌های n_1, n_2 از دو متغیر تصادفی برنولی در اختیار داشته باشیم، در این صورت برای مقایسه پارامترهای p_1, p_2 متغیرهای تصادفی از تفاضل این دو نسبت استفاده می‌کنیم به این ترتیب می‌توان نشان داد که یک فاصله اطمینان $(1-\alpha)$ درصدی برای تفاضل دو نسبت از دو متغیر تصادفی برنولی عبارتست از:

$$p \left(\bar{X}_1 - \bar{X}_2 - Z_{1-\frac{\alpha}{2}} \sqrt{\frac{\bar{X}_1(1-\bar{X}_1)}{n_1} + \frac{\bar{X}_2(1-\bar{X}_2)}{n_2}} \leq p_1 - p_2 \leq \bar{X}_1 - \bar{X}_2 + Z_{1-\frac{\alpha}{2}} \sqrt{\frac{\bar{X}_1(1-\bar{X}_1)}{n_1} + \frac{\bar{X}_2(1-\bar{X}_2)}{n_2}} \right) = 1-\alpha$$

$$\bar{X}_2 = \hat{P}_2, \bar{X}_1 = \hat{P}_1 \text{ که در آن}$$

۵۶ مثال ۲۳: در دو شهر A و B احتمال متولد شدن نوزادان پسر را با متغیرهای تصادفی برنولی X_1, X_2 نشان می‌دهیم. در شهر A از ۳۰۰

نوزاد متولد شده ۱۸۰ نوزاد پسر می‌باشند و در شهر B از ۴۰۰ نوزاد متولد شده ۱۹۰ نوزاد پسر متولد شده‌اند. می‌خواهیم احتمال متولد شدن نژاد

پسر را در دو شهر A و B مقایسه کنیم. برای این منظور از یک فاصله اطمینان ۹۵ درصدی استفاده می‌کنیم مطلوبست حدود فاصله اطمینان؟

حل: معلومات مساله عبارتند از:

$$n_1 = 300$$

$$\bar{X}_1 = \hat{p}_1 = \frac{180}{300} = \frac{18}{30} = 0.6$$

$$n_2 = 400$$

$$\bar{X}_2 = \hat{p}_2 = \frac{190}{400} = 0.475$$

$$1 - \alpha = 0.95 \rightarrow \alpha = 0.05 \rightarrow Z_{1-\frac{\alpha}{2}} = Z_{0.975} = 1.96$$

$$L_1 = \bar{X}_1 - \bar{X}_2 - Z_{1-\frac{\alpha}{2}} \sqrt{\frac{(1-\bar{X}_1)\bar{X}_1}{n_1} + \frac{\bar{X}_2(1-\bar{X}_2)}{n_2}}$$

$$= 0.6 - 0.475 - 1.96 \sqrt{\frac{(0.6) \cdot 0.4}{300} + \frac{(0.475) \cdot 0.525}{400}} = 0.05$$

$$L_2 = \bar{X}_1 - \bar{X}_2 + Z_{1-\frac{\alpha}{2}} \sqrt{\frac{\bar{X}_1(1-\bar{X}_1)}{n_1} + \frac{\bar{X}_2(1-\bar{X}_2)}{n_2}}$$

$$= 0.6 - 0.475 + 1.96 \sqrt{\frac{(0.6) \cdot 0.4}{300} + \frac{(0.475) \cdot 0.525}{400}} = 0.198$$

$$p(0.05 \leq p_1 - p_2 \leq 0.198) = 0.95$$

$$\Rightarrow p(p_2 + 0.05 \leq p_1 - p_2 + 0.198) = 0.95$$

بر اساس مساوی فوق احتمال متولد شدن نوزاد پسر در شهر A با اطمینان ۹۵ درصد در بازه $(p_2 + 0.05, p_2 + 0.198)$ قرار دارد.

فاصله اطمینان ۹۵٪

همچنین می‌توان نوشت:

S-۱۴

۵۷. ۱۰. ۶. فاصله اطمینان برای نسبت واریانسهای دو متغیر تصادفی نرمال

دو نمونه به حجم n_1 و n_2 را با واریانس نمونه‌ای S_1^2 و S_2^2 از دو جامعه نرمال با واریانس σ_1^2 و σ_2^2 استخراج می‌کنیم. برای مقایسه

واریانسهای دو جامعه از نسبت آنها $(\frac{S_1^2}{S_2^2})$ استفاده می‌کنیم. قبلاً نشان دادیم که متغیر تصادفی $\frac{S_1^2 \sigma_2^2}{S_2^2 \sigma_1^2}$ دارای توزیع F با $n_1 - 1$, $n_2 - 1$

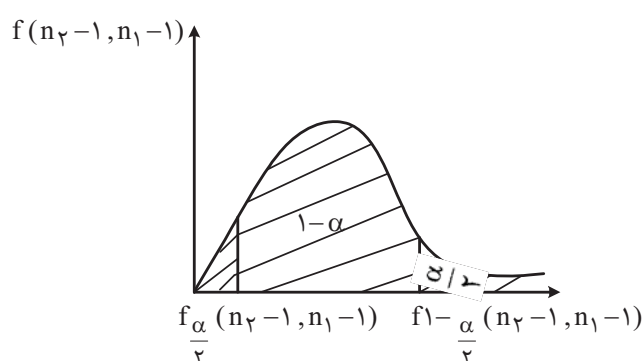
درجه آزادی است به این ترتیب می‌توان فاصله اطمینان $(1 - \alpha) 100$ درصدی را برای نسبت $\frac{\sigma_1^2}{\sigma_2^2}$ بصورت زیر بدست آورد:

$$p\left(\frac{S_1^2}{S_2^2} \frac{1}{f_{1-\frac{\alpha}{2}}(n_1-1, n_2-1)} \leq \frac{\sigma_1^2}{\sigma_2^2} \leq \frac{S_1^2}{S_2^2} \frac{1}{f_{1-\frac{\alpha}{2}}(n_1-1, n_2-1)}\right) = 1 - \alpha$$

با توجه به اینکه $f_{\frac{\alpha}{2}}(n_1, n_2) = \frac{1}{f_{1-\frac{\alpha}{2}}(n_2, n_1)}$ بنابراین رابطه فوق به صورت زیر ساده می‌شود:

$$p\left(\frac{S_1^2}{S_2^2} \frac{1}{f_{1-\frac{\alpha}{2}}(n_1-1, n_2-1)} \leq \frac{\sigma_1^2}{\sigma_2^2} \leq \frac{S_1^2}{S_2^2} f_{1-\frac{\alpha}{2}}(n_1-1, n_2-1)\right) = 1 - \alpha$$

که در آن $f_{\frac{\alpha}{2}}(n_1-1, n_2-1)$ مقداری از متغیر تصادفی F با n_1-1, n_2-1 درجه آزادی است بطوریکه سطح زیر منحنی به ازای مقادیر



کوچتر از آن برابر $\frac{\alpha}{2}$ باشد. به شکل زیر توجه کنید:

توجه کنید در تمام فواصل اطمینان که در روابط آنها از S^2 استفاده شده است اگر میانگین جامعه (μ) معلوم باشد به جای استفاده از

$$S^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \bar{X})^2 \quad \text{از} \quad S^2 = \frac{1}{n-1} \sum_{i=1}^n (x_i - \mu)^2 \quad \text{یعنی:}$$

مثال ۵۸: ۲۴: در یک کارخانه از دو ماشین A و B برای برش میله‌های فولادی استفاده می‌شود در صورتی که واریانس طول میله‌های بریده شده توسط هر یک از ماشینها سه برابر دیگری شود تصمیم به تعمیر ماشین مربوطه می‌گیریم. در یک نمونه ۲۵ تایی از میله‌های بریده شده توسط ماشین A مقدار S_1^2 برابر ۱۷۰ می‌باشد، همچنین در یک نمونه ۳۱ تایی از میله‌های بریده شده توسط B مقدار S_2^2 برابر ۲۵ می‌باشد. با اطمینان ۹۵٪ برای نسبت $\frac{\sigma_1^2}{\sigma_2^2}$ چه تصمیمی می‌توان در مورد ماشین A و B گرفت؟

حل: داریم:

$$n_1 = 25$$

$$n_2 = 31$$

$$S_1^2 = 170$$

$$S_2^2 = 25$$

$$1-\alpha = 0.95 \rightarrow \alpha = 0.05 \rightarrow f_{1-\frac{\alpha}{2}}(n_1-1, n_2-2) = f_{0.975}(24, 30) = 1/94$$

$$f_{\frac{\alpha}{2}}(n_1-1, n_2-2) = f_{0.025}(24, 30) = 1/89$$

به این ترتیب یک فاصله اطمینان ۹۵ درصدی عبارتست از:

$$p\left(\frac{170}{25} \frac{1}{1/94} \leq \frac{\sigma_1^2}{\sigma_2^2} \leq \frac{170}{25} \frac{1}{1/89}\right) = 0.95$$

$$\Rightarrow p\left(\frac{3}{5} \leq \frac{\sigma_1^2}{\sigma_2^2} \leq \frac{12}{85}\right) = 0.95$$

بنابراین نسبت واریانسها با ۹۵ درصد اطمینان در بازه $(3/5, 12/85)$ قرار دارد.

$$p(3/5, 12/85 \mid \sigma_1^2, \sigma_2^2) = 0.95$$

و در نتیجه:

یعنی با ۹۵ درصد اطمینان می‌توان گفت که واریانس ماشین A سه برابر واریانس ماشین B می‌باشد بنابراین ماشین A بایستی تعمیر شود.

۵۹. ۶. ۷ فاصله اطمینان برای تفاضل میانگین‌های دو جامعه نرمال

از دو جامعه نرمال با واریانسهای σ_1^2 و σ_2^2 دو نمونه به حجم n_1 و n_2 می‌گیریم. قبلاً نشان دادیم که $Z = \frac{\bar{X}_1 - \bar{X}_2 - (\mu_1 - \mu_2)}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}}$ یک

متغیر تصادفی نرمال استاندارد می‌باشد. بنابراین می‌توانیم یک فاصله اطمینان $(1-\alpha) \cdot 100$ درصد برای تفاضل میانگین‌های دو جامعه نرمال با واریانس معلوم بصورت زیر بدست آوریم:

$$p(\bar{X}_1 - \bar{X}_2 - Z_{1-\frac{\alpha}{2}} \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}} \leq \mu_1 - \mu_2 \leq \bar{X}_1 - \bar{X}_2 + Z_{1-\frac{\alpha}{2}} \sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}) = 1-\alpha$$

و در صورتی که واریانسهای دو جامعه نامعلوم اما مساوی باشند ابتدا واریانس مشترک دو نمونه را بصورت زیر برآورد می‌کنیم:

$$\sigma_1^2 = \sigma_2^2 = \sigma^2$$

$$\hat{\sigma}^2 = S_p^2 = \frac{(n_1-1) S_1^2 + (n_2-1) S_2^2}{n_1 + n_2 - 2}$$

فاصله اطمینان در این حالت از رابطه زیر بدست می‌آید:

$$p(\bar{X}_1 - \bar{X}_2 - t_{1-\frac{\alpha}{2}}(n_1+n_2-2) \sqrt{\frac{1}{n_1} + \frac{1}{n_2}} \leq \mu_1 - \mu_2 \leq \bar{X}_1 - \bar{X}_2 + t_{1-\frac{\alpha}{2}}(n_1+n_2-2) S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}) = 1-\alpha$$

۶۰ مثال ۲۵: میزان برق مصرفی دو شهر A و B یک متغیر تصادفی نرمال می‌باشد برای مقایسه متوسط مصرف یک نمونه ۱۵ تایی از دو شهر

انتخاب می‌کنیم و نتایج زیر را بدست می‌آوریم:

$$\bar{X}_A = ۲۵۹۰ \quad \sigma_A^2 = ۱۵۰$$

$$\bar{X}_B = ۲۵۶۰ \quad \sigma_B^2 = ۱۰۰$$

مطلوبست:

الف) با فاصله اطمینان ۹۹ درصدی چه نظری می‌توان درباره تفاضل متوسط مصرف در دو شهر داشت؟

ب) اگر برای دو نمونه داشته باشیم $S_A^2 = ۱۶۵$, $S_B^2 = ۱۲۰$ یک فاصله اطمینان ۹۹ درصدی برای حالتی که واریانسها نامعلوم باشند و مساوی بسازید؟

حل: الف)

$$1-\alpha = ۰/۹۹ \rightarrow \alpha = ۰/۰۱ \rightarrow Z_{1-\frac{\alpha}{2}} = Z_{۰/۹۹} = ۲/۵۷۵$$

$$\rightarrow L_1 = ۲۵۹۰ - ۲۵۶۰ - ۲/۵۷۵ \sqrt{\frac{۱۵۰}{۱۵} + \frac{۱۰۰}{۱۵}} = ۱۹/۴۸$$

$$L_2 = ۲۵۹۰ - ۲۵۶۰ - ۲/۵۷۵ \sqrt{\frac{۱۵۰}{۱۵} + \frac{۱۰۰}{۱۵}} = ۴۰/۵۱$$

بنابراین با ۹۹ درصد اطمینان می‌توانیم بگوییم که $۱۹/۴۸ \leq \mu_1 - \mu_2 \leq ۴۰/۵۱$ و چون $\mu_1 - \mu_2$ در یک بازه مثبت قرار گرفته‌اند می‌توانیم با ۹۹ درصد اطمینان بگوییم که متوسط مصرف شهر A از شهر B بیشتر است.

ب) ۶۱

$$S_p^2 = \frac{(n_1-1) S_1^2 + (n_2-1) S_2^2}{n_1 + n_2 - 2} = \frac{۱۴ (۱۲۰ + ۱۶۵)}{۱۵ + ۱۵ - 2} = ۱۴۲/۵ \Rightarrow S_p = ۱۱/۹۳$$

$$1-\alpha = 0.99 \rightarrow \alpha = 0.01 \rightarrow t_{1-\frac{\alpha}{2}}(n_1+n_2-2) = t_{0.995}(28) = 2.76$$

$$L_1 = \bar{X}_1 - \bar{X}_2 - t_{1-\frac{\alpha}{2}}(n_1+n_2-2) S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}} = 2590 - 2560 - 2.76(11/93) \sqrt{\frac{1}{15} + \frac{1}{15}} = 17.97$$

$$L_2 = \bar{X}_1 - \bar{X}_2 - t_{1-\frac{\alpha}{2}}(n_1+n_2-2) S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}} = 2590 - 2560 + 2.76(11/93) \sqrt{\frac{1}{15} + \frac{1}{15}} = 42$$

ملاحظه می کنید که در حالت دوم که واریانس مشترک را بصورت تقریبی محاسبه نمودیم باز هم فاصله اطمینان مقداری نزدیک به حالت الف شده است.

فصل یازدهم

۱۱-۱ آزمون فرض

در فصل قبل با روشهای برآورد پارامترهای مجهول جامعه آشنا شدید، به عبارتی نشان دادیم که چگونه می‌توان با استفاده از نتایج حاصل از نمونه‌گیری پارامترهای مجهول یک جامعه آماری با توزیع معلوم را بدست آورد. حال در این فصل نشان می‌دهیم که چگونه می‌توان در مورد صحت یک فرضیه بر اساس یک جامعه آماری قضاوت نمود. به عنوان مثال فرض می‌کنیم که یک داروی بخصوص بر روی یک بیماری اثر مثبتی دارد به این ترتیب می‌بایستی برای اثبات این فرضیه، عده‌ای از بیماران را انتخاب کرده و دارو را روی آنها آزمایش می‌کنیم و با توجه به نتایج بدست آمده در مورد صحت یا عدم صحت فرضیه مورد نظر تصمیم می‌گیریم. روشی که در آن فرض مربوطه را پذیرفته یا رد می‌کنیم به آزمون فرض معروف می‌باشد، در این فصل نحوه آزمون فرضهای مختلف را معرفی می‌کنیم.

۱۱-۲

۱.۱ فرضها و آزمونها

تعریف دقیق فرض و آزمون فرض عبارتند از:

فرض: یک فرض عبارتست از گزاره‌ای درباره قانون احتمال یک متغیر تصادفی (که قابل نمونه‌گیری باشد)
آزمون فرض: نمونه‌گیری از متغیر تصادفی مربوطه و تصمیم‌گیری در مورد پذیرفتن یا رد فرض مورد نظر بر اساس نمونه بدست آمده.
بطور کلی نتیجه یک آزمون فرض یکی از چهار مورد زیر می‌باشد:

- ۱- فرض را رد می‌کنیم هنگامی که واقعاً صحیح باشد.
 - ۲- فرض را رد می‌کنیم هنگامی که واقعاً صحیح نباشد.
 - ۳- فرض را رد می‌پذیریم هنگامی که واقعاً صحیح باشد.
 - ۴- فرض را رد می‌پذیریم هنگامی که واقعاً صحیح نباشد.
- برای روشن شدن مطلب به مثال زیر توجه کنید:

۱۱-۳ مثال ۱: گروهی از پزشکان معتقدند یک دارو نیروزا می‌باشد. برای تصمیم‌گیری در این مورد آنرا روی ۲۰ نفر از ورزشکاران آزمایش می‌کنند. در صورتی که دارو در بیشتر از ۷۰ درصد مواقع مثبتی بر عملکرد عادی ورزشکاران داشته باشد، پزشکان آنرا جزء داروهای نیروزا طبقه بندی می‌کنند. متغیر تصادفی X را برابر تعداد ورزشکارانی در بین ۲۰ نفر در نظر می‌گیریم که دارو تاثیر مثبتی بر عملکرد آنها داشته است. در این صورت X یک متغیر تصادفی دو جمله‌ای با پارامترهای $B(p, 20)$ می‌باشد. که در آن p درصد موثر بودن دارو می‌باشد و مقداری نامعلوم می‌باشد. نیروزا بودن دارو به این معنی است که $p > 0.70$ به عبارتی اگر از روی نمونه‌ها به این نتیجه برسیم که $p > 0.70$ می‌باشد در این صورت می‌توان گفت که دارو نیروزا بوده است. نتایج حاصل از نمونه‌ها یکی از دو حالت (فرض) زیر را نتیجه می‌دهد:

- ۱- داروی مورد نظر نیروزا می‌باشد: $p > 0.70$
- ۲- داروی مورد نظر نیروزا نمی‌باشد: $p > 0.70$

۱۱-۴ در روش آزمون یکی از فرضها را فرض خنثی یا فرض صفر در نظر می‌گیریم و آنرا با H_0 نمایش می‌دهیم به عبارتی H_0 در این مثال همان فرض است که مورد ادعای پزشکان می‌باشد یعنی فرض داروی مورد نظر نیروزا می‌باشد $(p > 0.70)$ فرض H_0 خواهد بود. در مقابل هر فرضی که فرض H_0 را رد کند فرض مخالف یا مکمل نامیده می‌شود و با H_1 نمایش داده می‌شود. که در این مثال فرض داروی مورد نظر نیروزا نمی‌باشد $(p \leq 0.70)$ فرض H_1 می‌باشد. بنابراین:

دارو نیروزا است $H_0 = p > 0.70$

دارو نیروزا نیست $H_1 = p \leq 0.70$

توجه کنید که فرض مخالف صورتهای متفاوتی می تواند داشته باشد مثلاً اگر فرض H_0 برابر $p = 0.7$ می بود در این صورت فرض H_1 هر یک از حالات زیر می توانست باشد:

$$H_1 : p > 0.7$$

$$H_1 : p \neq 0.7$$

$$H_1 : p = 0.5$$

تصمیم گیری در مورد اینکه کدام فرض را فرض مخالف در نظر بگیریم اهمیت زیادی دارد و معمولاً به صورت مساله بستگی دارد. در ایم مثال پزشکان برای اثبات ادعای خود ۲۰ ورزشکار را انتخاب نموده و با توجه به متغیر تصادفی X این احتمال وجود دارد که نتایج طوری بدست آیند که پزشکان دارو را نیروزا معرفی کنند در حالی که واقعاً دارو نیروزا بوده است.

۱۱-۵ به عنوان مثالی دیگر فرض کنید یک سکه را ۲۰ مرتبه پرتاب کنیم و ۱۶ شیر مشاهده کنیم بوضوح مشاهده ۱۶ شیر تنها با ۲۰ مرتبه پرتاب سکه دلیلی بر ناسالم بودن سکه نمی باشد در مثال داروی نیروزا هم این مطلب صدق می کند. بنابراین در حالت کلی چهار حالت زیر بدست می آید:

۱- H_0 را می پذیریم هنگامی که واقعاً صحیح باشد.

۲- H_0 را رد می کنیم هنگامی که واقعاً صحیح باشد.

۳- H_0 را می پذیریم هنگامی که واقعاً صحیح نباشد.

۴- H_0 را رد می کنیم هنگامی که واقعاً صحیح نباشد.

مشاهده می کنید که حالت دوم و سوم نشان دهنده خطا در تصمیم گیری می باشند حالت دوم را خطای نوع اول (ریسک فروشنده می نامیم و حالت سوم را خطای نوع دوم (ریسک مشتری) می نامیم. خطای نوع اول را با α نمایش می دهند و آنرا سطح معنی دار یا سطح تشخیص آزمون می نامند و خطای نوع دوم را با B نمایش می دهند.

۱۱-۱۱. ۲ ناحیه بحرانی و آماره آزمون

برای بیان ناحیه بحرانی و آماره آزمون مجدداً به مثال ۱ توجه کنید. در مثال ۱ نشان دادیم که:

$$H_0 : p > 0.70$$

$$H_1 : p \leq 0.70$$

چگونگی انجام آزمون به این ترتیب است که متغیر تصادفی X را برابر تعداد ورزشکاران که دارو بروی عملکرد آنها تاثیر مثبتی داشته در نظر می گیریم

در این صورت اگر مقدار مشاهده شده X مثلاً برابر ۱۵ باشد در نتیجه $p = \frac{15}{20} = \frac{3}{4} = 0.75$ و فرض H_0 را قبول می کنیم زیرا $p > 0.70$ اما

اگر مقدار مشاهده شده X مثلاً برابر ۵ باشد بوضوح فرض H_0 را رد می کنیم و فرض H_1 را می پذیریم زیرا $p = \frac{5}{20} = 0.25$ و $p \leq 0.70$.

در حالت کلی زمانی فرض H_0 را می پذیریم که مقدار مشاهده شده X از عدد ۱۴ بیشتر باشد زیرا:

$$p = \frac{14}{20} = 0.70$$

و اگر مقدار مشاهده شده X کوچکتر یا مساوی ۱۴ باشد فرض H_0 را رد می کنیم یعنی اگر مقادیر مشاهده شده X متعلق به مجموعه

$$C = \{x \mid x \leq 14\}$$

H_0 باشد آنگاه H_0 را رد می کنیم و در غیر این صورت H_0 را می پذیریم. به آماره X که بر اساس مقادیر مشاهده شده آن

H_0 را رد می کنیم یا می پذیریم آماره آزمون گویند و به ناحیه C که به ازای مقادیر آن H_0 رد می شود ناحیه بحرانی می گوییم. به تعریف زیر توجه

کنید:

۷-۱۱ تعریف: به آماره $T = T(X_1, X_2, \dots, X_n)$ که بر اساس مقادیر مشاهده شده آن یک فرض را می‌پذیریم یا رد می‌کنیم، آماره آزمون گویند. و به ازای مجموعه مقادیری از این آماره که بر اساس آنها فرض H_0 رد می‌شود ناحیه بحرانی آزمون می‌گوییم و با نماد C نمایش می‌دهیم و به متمم ناحیه بحرانی که با C' نمایش می‌دهیم ناحیه پذیرش آزمون می‌گوییم.

با توجه به تعریف در آزمون یک فرض به این صورت عمل می‌کنیم که ابتدا مقادیر نمونه‌های $X_1, X_2, \dots, X_n$ را جمع آوری می‌کنیم و بر اساس آنها آماره آزمون را که بشکل $t = T(x_1, x_2, \dots, x_n)$ می‌باشد تشکیل می‌دهیم حال اگر $t \in C$ باشد فرض H_0 را رد می‌کنیم و در غیر این صورت آن را می‌پذیریم.

۸.۱۱. ۳ خطای آزمون

اشاره کردیم که رد تصمیم‌گیری ممکن است دو نوع خطای نوع اول (α) و نوع دوم (β) رخ دهد، که در هر دو حالت نتیجه مطلوب نمی‌باشد و در عمل می‌خواهیم تا جای ممکن مقدار خطای α و β را کاهش دهیم. با استفاده از احتمالات شرطی خطای نوع اول و دوم بصورت زیر بدست می‌آیند:

$$\alpha = p(H_0 \text{ درست باشد} | T(X_1, \dots, X_n) \in C) = p(T(X_1, \dots, X_n) \in C | H_0 \text{ درست باشد})$$

$$\beta = p(H_1 \text{ درست باشد} | T(X_1, \dots, X_n) \notin C) = p(T(X_1, \dots, X_n) \notin C | H_1 \text{ درست باشد})$$

اگر در آزمون فرض به این نتیجه برسیم که H_0 باید رد شود و در حالی که H_1 واقعاً درست باشد یعنی مقدار احتمال زیر را داشته باشیم:

$$p(H_1 \text{ درست باشد} | H_0 \text{ رد شود}) = p(H_0 \text{ واقعاً نادرست باشد} | H_0 \text{ رد شود}) = p$$

آشکار است که هر چه مقدار این احتمال بیشتر باشد آزمون نتیجه دقیق‌تری بدست داده است به همین دلیل به این احتمال توان آزمون گفته می‌شود. و با علامت β^* نمایش داده می‌شود.

توجه کنید که مقدار β^* برابر با $1 - \beta$ می‌باشد:

$$\beta^* = p(H_1 \text{ درست باشد} | H_0 \text{ پذیرفته شود}) = 1 - p(H_1 \text{ درست باشد} | H_0 \text{ رد شود}) = 1 - \beta$$

۹ مثال ۲: شخصی ادعا می‌کند که سکه استفاده شده برای تعیین زمین بازی در یک مسابقه فوتبال یکنواخت نبوده و با احتمال 0.6 شیر می‌آمده

است. برای تحقیق صحت ادعای وی سکه را ۲۰ مرتبه پرتاب می‌کنیم مطلوبست:

الف) تعیین ناحیه بحرانی آزمون و فرضهای آزمون.

ب) محاسبه مقدار احتمال خطای نوع اول و خطای نوع دوم.

ج) محاسبه توان آزمون.

حل: الف) X را تعداد شیرهای بدست آمده در ۲۰ مرتبه پرتاب سکه در نظر می‌گیریم در این صورت X یک متغیر تصادفی دو جمله‌ای با پارامتر p و ۲۰ می‌باشد. $X \sim \beta(20, p)$.

طبق تعریف فرض صفر خلاف ادعای مطرح شده می‌باشد بنابراین:

$$\begin{cases} H_0 : p = \frac{1}{2} \\ H_1 : p = \frac{6}{10} \end{cases}$$

با توجه به اینکه در متغیر تصادفی برنولی $\mu = np$ می‌باشد بنابراین اگر در ۲۰ مرتبه پرتاب سکه حداقل تعداد $12 = 20 \times \frac{6}{10}$ شیر مشاهده کنیم

می‌بایستی فرض H_0 را رد کنیم به این ترتیب ناحیه بحرانی عبارتست از:

$$C = \{x | x \geq 12\}$$

ب) محاسبه خطای نوع اول:

$$\alpha = p(X \in C \mid \text{درست باشد}) = p(X \geq 12 \mid p = \frac{1}{4})$$

$$1 - p(X \leq 11 \mid p = \frac{1}{4}) = 1 - 0.7483 = 0.2517$$

محاسبه خطای نوع دوم:

$$\beta = p(X \notin C \mid \text{درست } H_1) = p(X < 12 \mid p = 0.6)$$

$$= p(X \leq 11 \mid p = 0.6) = 0.4044$$

ج) محاسبه توان آزمون:

$$\beta^* = 1 - \beta = 0.5956$$

ملاحظه می کنید که احتمال خطای نوع اول کم و احتمال خطای نوع دوم زیاد می باشد.

۱۰-۱۱ مثال ۳: در مثال قبل اگر ناحیه بحرانی بصورت $C = \{x \mid x \geq 13\}$ باشد احتمال خطای نوع اول و دوم و توان آزمون را محاسبه کنید؟
حل:

$$\alpha = p(X \geq 13 \mid p = \frac{1}{4}) = 1 - p(X \leq 12 \mid p = \frac{1}{4}) = 1 - 0.8684 = 0.1316$$

$$\beta = p(X < 13 \mid p = 0.6) = p(X \leq 12 \mid p = 0.6) = 0.5841$$

$$\beta^* = 1 - \beta = 0.4159$$

ملاحظه می کنید که در حالت دوم با متغیر ناحیه بحرانی مقدار احتمال نوع اول کاهش یافت اما این مساله موجب افزایش احتمال خطای نوع دوم شده است.

بنابراین می بایستی ناحیه بحرانی طوری انتخاب شود که همزمان با در نظر گرفتن یک مقدار حداکثر برای خطای نوع اول باعث حداقل نمودن نوع دوم شود. به این ترتیب توان آزمون نیز حداکثر می شود.

در مثال بعد با قرار دادن یک مقدار معین برای احتمال خطای نوع اول (α) تلاش می نیم احتمال خطای نوع دوم را تا جای ممکن کاهش دهیم و در نتیجه توان آزمون را افزایش دهیم.

۱-۱۱-۱۱ مثال ۴: X یک متغیر تصادفی نرمال با میانگین مجهول μ و واریانس σ^2 می باشد. آزمون فرض زیر را در نظر بگیرید:

$$\begin{cases} H_0 : \mu = 0 \\ H_1 : \mu = 1 \end{cases}$$

اگر از X یک نمونه ۲۵ تایی بگیریم و ناحیه بحرانی بصورت $C = \{ (X_1, \dots, X_{25}) \mid \bar{X} \geq C \}$ باشد مقدار C را طوری بدست بیاورید که $\alpha = 0.1$ باشد و سپس احتمال خطای نوع دوم و توان آزمون را محاسبه کنید؟

حل: ابتدا توجه کنید که $\bar{X}$ دارای میانگین μ و واریانس $\frac{\sigma^2}{n}$ می باشد. یعنی $\bar{X} \sim (\mu, \frac{\sigma^2}{n})$ به این ترتیب:

$$\alpha = 0.1 = p(\bar{X} > C \mid \mu = 0) = p\left(\frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} > \frac{C - \mu}{\frac{\sigma}{\sqrt{n}}} \mid \mu = 0\right)$$

$$= p\left(\frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} > \frac{C - 0}{\frac{\sigma}{\sqrt{n}}} \right) = p(Z > 1.66 C) = 0.1$$

$$\Rightarrow 1 - p(Z \leq 1.66 C) = 0.1 \Rightarrow p(Z \leq 1.66 C) = 0.9$$

رابطه اخیر معادل است با اینکه $C = 1.66 Z_{0.9}$ از جدول مقدار $Z_{0.9}$ برابر است با ۱.۲۸ بنابراین:

$$Z_{./9} = 1/28 = 1/66 C \Rightarrow C = 0/77$$

به این ترتیب ناحیه بحرانی بصورت زیر بدست می‌آید:

$$C = \{ (X_1, \dots, X_{25}) \mid \bar{X} > 0/77 \}$$

ملاحظه می‌کنید که در این مثال برای اینکه بتوانیم مقدار خطای نوع اول را به دلخواه کاهش دهیم به ناچار می‌بایستی بازه ناحیه بحرانی را متغیر فرض می‌کردیم، حال مقدار خطای نوع دوم را بدست می‌آوریم:

$$\beta = p(\bar{X} \leq C \mid \mu = 1) = p\left(\frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} \leq \frac{0/77 - \mu}{\frac{\sigma}{\sqrt{n}}} \mid \mu = 1\right)$$

$$= p\left(\frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} \leq \frac{0/77 - 1}{\frac{3}{5}}\right) = p(Z \leq -0/38) = N_Z(-0/38) = 0/3520$$

$$\beta^* = 1 - \beta = 1 - 0/3520 = 0/648$$

توان آزمون:

۱۲-۱۱. ۴ انواع فرضیه

فرض $\mu = 2$: H_0 را در نظر بگیرید، این فرض با بیان اینکه میانگین جامعه برابر ۲ می‌باشد توزیع جامعه را معلوم می‌کند و نشان می‌دهد که پارامتر مجهول جامعه میانگین برابر عدد ۲ می‌باشد، به این نوع فرضیه که توزیع جامعه را کاملاً مشخص می‌سازند فرضیه‌های ساده می‌گوییم. حال فرض $\mu \geq 2$: H_1 را در نظر بگیرید، واضح است که اگر این فرض درست باشد با بیان یک بازه برای پارامتر مجهول جامعه، توزیع را بصورت دقیق مشخص نمی‌کند به این فرضیه که با درست بودنشان پارامتر مجهول جامعه و در نتیجه توزیع جامعه بصورت دقیق مشخص نمی‌شود فرضیه‌های مرکب می‌گوییم.

به عنوان مثال در مثال ۴ فرضیه‌های $H_0: \mu = 0$ و $H_1: \mu = 1$ فرضیه‌های ساده می‌باشند اما در مثال ۱ فرضیه‌های $H_0: p \geq 0/6$ و $H_1: p < 0/6$ فرضیه‌های مرکب می‌باشند.

۱۳-۱۱. ۵ انواع آزمونها

اگر θ یک پارامتر نامعلوم جامعه باشد و بخواهیم آزمونهایی در مورد این پارامتر انجام دهیم، آزمون هر فرض آماری که در آن فرضیه مقابل (H_1) یک طرفه باشد آزمون یک طرفه نامیده می‌شود. به عنوان مثال برای مقدار ثابت θ_0 از θ آزمونهایی زیر همگی یک طرفه می‌باشند:

$$\begin{cases} H_0: \theta = \theta_0 \\ H_1: \theta \geq \theta_0 \end{cases} \quad \begin{cases} H_0: \theta \geq \theta_0 \\ H_1: \theta < \theta_0 \end{cases}$$

اگر در یک آزمون فرضیه مقابل (H_1) دو طرفه باشد آن آزمون را آزمون دو طرفه می‌نامیم. مانند آزمون زیر:

$$\begin{cases} H_0: \theta = \theta_0 \\ H_1: \theta \neq \theta_0 \end{cases}$$

۱۴-۱۱. ۶ مراحل انجام یک آزمون فرض

برای انجام یک آزمون فرض می‌بایستی مراحل زیر را بصورت گام به گام طی نمود:

۱- تعیین فرضیه‌های صفر H_0 و فرض مقابل H_1 .

۲- تعیین یک سطح معنی‌دار α که معمولاً یکی از اعداد ۰/۰۱، ۰/۰۵، یا ۰/۱ در نظر گرفته می‌شود.

۳- تعیین آماره آزمون $T = T(X_1, \dots, X_n)$ که عموماً بر اساس برآوردگر نقطه‌ای پارامتر مجهول θ بدست می‌آید.

۴- تعیین ناحیه بحرانی آزمون که از روی آماره آزمون، فرض مقابل آزمون و سطح معنی‌دار α بدست می‌آید.

۵- محاسبه مقدار آماره آزمون که از روی نمونه‌های تصادفی $X_1, X_2, \dots, X_n$ بدست می‌آید.

۶- نتیجه‌گیری - اگر مقدار محاسبه شده آماره آزمون درون ناحیه بحرانی باشد فرض H_0 را رد و در غیر این صورت فرض H_0 را می‌پذیریم. توجه کنید که عموماً در حل مسایل آزمون فرض، با در نظر گرفتن انواع آزمون فرضها و تعیین آماره آزمون و ناحیه بحرانی برای هر یک از آزمون فرضها محاسبه مراحل ۳ و ۴ بسیار ساده‌تر می‌شود.

۱۵-۱۱.۷ آزمون فرضهای آماری برای پارامترهای جامعه

برای سادگی انجام آزمونهای فرض بهترین راه محاسبه آماره آزمون و ناحیه بحرانی برای فرضهای متفاوت می‌باشد. بنابراین در این بخش ابتدا آزمون فرض را روی میانگین یک جمعیت را زمانی که واریانس معلوم می‌باشد انجام می‌دهیم:

فرض می‌کنیم X یک متغیر تصادفی با میانگین مجهول μ و واریانس معلوم σ^2 باشد. نمونه‌های تصادفی $X_1, X_2, \dots, X_n$ را انتخاب می‌کنیم، می‌خواهیم آزمونهایی را روی میانگین μ انجام دهیم. در این صورت سه حالت کلی زیر را در نظر می‌گیریم:

$$1- \text{آزمون فرض} \begin{cases} H_0 : \mu = \mu_0 \\ H_1 : \mu > \mu_0 \end{cases} \quad \text{که در آن } H_0 \text{ رد می‌شود اگر } \bar{X} > C$$

$$2- \text{آزمون فرض} \begin{cases} H_0 : \mu = \mu_0 \\ H_1 : \mu < \mu_0 \end{cases} \quad \text{که در آن } H_0 \text{ رد می‌شود اگر } \bar{X} > C$$

$$3- \text{آزمون فرض} \begin{cases} H_0 : \mu = \mu_0 \\ H_1 : \mu \neq \mu_0 \end{cases} \quad \text{که در آن } H_0 \text{ رد می‌شود اگر } \bar{X} < C_1 \text{ یا } \bar{X} > C_2$$

توجه کنید که در هر حالت مقدار C که بیانگر ناحیه بحرانی می‌باشد با توجه به سطح معنی‌دار α بدست می‌آید. حال برای هر حالت آماره آزمون و مقدار C را محاسبه می‌کنیم:

۱۶-۱۱ حالت اول: H_0 رد می‌شود اگر $\bar{X} > C$ بنابراین خطای نوع اول عبارتست از:

اگر X یک متغیر تصادفی نرمال باشد یا $n \geq 30$ باشد قبلاً نشان دادیم که $Z = \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}}$ یک متغیر تصادفی نرمال استاندارد می‌باشد بنابراین:

$$Z = \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} \sim N(0, 1)$$

$$\alpha = P(\bar{X} > C \mid \mu = \mu_0) = P\left(\frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} > \frac{C - \mu_0}{\frac{\sigma}{\sqrt{n}}}\right) = P\left(Z > \frac{C - \mu_0}{\frac{\sigma}{\sqrt{n}}}\right)$$

$$\Rightarrow 1 - P\left(Z > \frac{C - \mu_0}{\frac{\sigma}{\sqrt{n}}}\right) = \alpha \Rightarrow P\left(Z > \frac{C - \mu_0}{\frac{\sigma}{\sqrt{n}}}\right) = 1 - \alpha$$

$$\Rightarrow \frac{C - \mu_0}{\frac{\sigma}{\sqrt{n}}} = Z_{1-\alpha} \Rightarrow C = \frac{\sigma}{\sqrt{n}} (Z_{1-\alpha}) + \mu_0$$

حال با توجه به مقدار C ناحیه بحرانی آزمون بصورت زیر بدست می‌آید:

$$\bar{X} > C \rightarrow \bar{X} > \mu_0 + \frac{\sigma}{\sqrt{n}} Z_{1-\alpha} \quad \text{یا} \quad \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} > Z_{1-\alpha}$$

نتیجه می‌گیریم که:

در آزمون $\begin{cases} H_0 : \mu = \mu_0 \\ H_1 : \mu > \mu_0 \end{cases}$ فرض H_0 رد می‌شود اگر و فقط اگر:

$$Z = \frac{\bar{X} - \mu_0}{\frac{\sigma}{\sqrt{n}}} > Z_{1-\alpha}$$

توجه کنید که اگر فرض H_0 بصورت $H_0 : \mu \leq \mu_0$ باشد می‌توان نشان داد که ناحیه بحرانی باز هم به صورت رابطه بالا می‌باشد.

۱۷-۱۱ حالت دوم: آزمون فرض $\begin{cases} H_0 : \mu = \mu_0 \\ H_1 : \mu < \mu_0 \end{cases}$ را رد می‌کنیم اگر $\bar{X} < C$ مشابه حالت قبل اگر C را محاسبه کنیم مقدار ناحیه بحرانی

بصورت زیر بدست می‌آید:

$$C = \mu_0 - Z_{1-\alpha} \frac{\sigma}{\sqrt{n}} \Rightarrow \bar{X} < C \rightarrow \bar{X} < \mu_0 - \frac{\sigma}{\sqrt{n}} Z_{1-\alpha}$$

$$Z = \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} < -Z_{1-\alpha}$$

بنابراین در حالت کلی در آزمون فرض $\begin{cases} H_0 : \mu = \mu_0 \\ H_1 : \mu < \mu_0 \end{cases}$ رد می‌شود اگر و فقط اگر $Z = \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} < -Z_{1-\alpha}$ در این حالت هم اگر فرض

H_0 بصورت $H_0 : \mu \geq \mu_0$ می‌بود باز هم ناحیه بحرانی بصورت رابطه فوق بدست می‌آمد.

۱۸-۱۱ حالت سوم: آزمون فرض $\begin{cases} H_0 : \mu = \mu_0 \\ H_1 : \mu \neq \mu_0 \end{cases}$ که در آن H_0 رد می‌شود اگر $\bar{X} > C_1$ یا $\bar{X} < C_2$ باشد بنابراین خطای نوع اول عبارتست

از:

$$\alpha = p(\bar{X} < C_1 \text{ یا } \bar{X} > C_2 \mid \mu = \mu_0) \rightarrow 1 - \alpha = p(C_1 < \bar{x} < C_2 \mid \mu = \mu_0)$$

$$p = \left(\frac{C_1 - \mu_0}{\frac{\sigma}{\sqrt{n}}} < \frac{\bar{X} - \mu}{\frac{\sigma}{\sqrt{n}}} < \frac{C_2 - \mu_0}{\frac{\sigma}{\sqrt{n}}} \right)$$

$$= p \left(\frac{C_1 - \mu_0}{\frac{\sigma}{\sqrt{n}}} < Z < \frac{C_2 - \mu_0}{\frac{\sigma}{\sqrt{n}}} \right) = 1 - \alpha$$

$$\frac{C_2 - \mu_0}{\frac{\sigma}{\sqrt{n}}} = - \frac{C_1 - \mu_0}{\frac{\sigma}{\sqrt{n}}} < Z_{1-\alpha}$$

که در آن:

و در نتیجه بدست می‌آوریم:

$$C_1 = \mu_0 - Z_{1-\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}$$

$$C_2 = \mu_0 + Z_{1-\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}$$

و ناحیه بحرانی به صورت زیر بدست می‌آید:

$$\bar{X} > \mu_0 + Z_{1-\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}} \quad \text{یا} \quad \bar{X} < \mu_0 - Z_{1-\frac{\alpha}{2}} \frac{\sigma}{\sqrt{n}}$$

که معادل است با:

$$\left| \frac{X' - \mu_0}{\frac{\sigma}{\sqrt{n}}} \right| > Z_{1-\frac{\alpha}{2}}$$

بنابراین در حالت کلی در آزمون فرض $\begin{cases} H_0: \mu = \mu_0 \\ H_1: \mu \neq \mu_0 \end{cases}$ فرض H_0 رد می‌شود اگر فقط اگر:

$$|Z| = \left| \frac{X' - \mu_0}{\frac{\sigma}{\sqrt{n}}} \right| > Z_{1-\frac{\alpha}{2}}$$

در مثالهای بعدی هر حالت را مورد بررسی قرار می‌دهیم.

۱۹-۱۱ مثال ۵: توان شکنندگی (مقدار نیروی لازم برای شکستن) کابلهایی که توسط یک شرکت تولید می‌شوند دارای میانگین ۱۸۰۰ پوند و انحراف معیار ۱۰۰ می‌باشند. محققان شرکت با اعمال تکنیک جدیدی که در مراحل ساخت اعمال کرده‌اند ادعا کرده‌اند که توان شکنندگی افزایش یافته است. برای آزمودن این ادعا یک نمونه ۵۰ تایی از کابلها تحت آزمون قرار می‌گیرند و میانگین توان شکنندگی ۱۸۵۰ پوند بدست می‌آید. آیا در سطح معنی‌دار ۰/۰۱ این ادعا پذیرفته است؟

حل: ابتدا هر یک از فرضهای صفر و مقابل را تعریف می‌کنیم:

$H_0: \mu = 1800$ هیچ تغییری در توان شکنندگی رخ نداده است.

$H_1: \mu > 1800$ توان شکنندگی افزایش یافته است.

ملاحظه می‌کنید که حالت اول رخ داده است و با توجه به مطالب ارایه شده آماره آزمون و ناحیه بحرانی بصورت زیر خواهند بود:

$$Z = \frac{X' - \mu_0}{\frac{\sigma}{\sqrt{n}}} > Z_{1-\frac{\alpha}{2}}$$

که در آن:

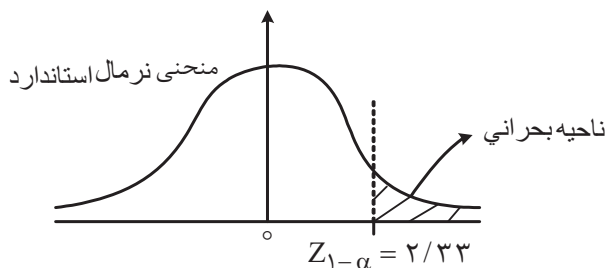
$$\alpha = 0.01$$

$$\Rightarrow 1 - \alpha = 1 - 0.01 = 0.99 \quad \Rightarrow \quad Z_{1-\alpha} = Z_{0.99} = 2.33$$

یعنی ناحیه بحرانی بصورت $Z > 2.33$ می‌باشد و اگر مقدار آماره Z بزرگتر از ۲/۳۳ باشد فرض H_0 را رد می‌کنیم و به عبارتی ادعای محققان را می‌پذیریم. مقدار آماره Z برابر است با:

$$Z = \frac{X' - \mu_0}{\frac{\sigma}{\sqrt{n}}} = \frac{1850 - 1800}{\frac{100}{\sqrt{50}}} = 3.55$$

بنابراین در سطح معنی‌دار $\alpha = 0.01$ نتایج نشان می‌دهند که $Z > 2.33$ می‌باشد و در نتیجه ادعای محققان را در مورد افزایش توان شکنندگی کابلها می‌پذیریم. در نمودار زیر ناحیه رد فرض H_0 را ملاحظه می‌کنید:



۲۰-۱۱ مثال ۶: عمر متوسط ۱۰۰ عدد از لامپهای مهتابی تولید شده توسط یک کارخانه برابر ۱۵۷۰ ساعت با انحراف معیار ۱۲۰ ساعت بدست آمده است. کارخانه تولید کننده ادعا می کند که عمر متوسط لامپها برابر ۱۶۰۰ ساعت می باشد در حالی که مصرف کنندگان این ادعا را قبول ندارند. صحت ادعای کارخانه سازنده را در سطح $\alpha = 0.05$ و $\alpha = 0.01$ بررسی کنید؟

حل: در این حالت دو فرض زیر را پیش رد داریم:

$$\begin{cases} H_0 : \mu = 1600 \\ H_1 : \mu \neq 1600 \end{cases}$$

ملاحظه می کنید که در این حالت یک آزمون دو طرفه انجام می گیرد که در آن آماره آزمون و ناحیه بحرانی بصورت زیر می باشد:

$$|Z| = \left| \frac{X' - \mu_0}{\frac{\sigma}{\sqrt{n}}} \right| > Z_{1-\frac{\alpha}{2}}$$

الف) ابتدا آزمون فرض را در سطح معنی دار $\alpha = 0.05$ انجام می دهیم:

$$\alpha = 0.05 \rightarrow 1 - \frac{\alpha}{2} = 0.975$$

$$\rightarrow Z_{1-\frac{\alpha}{2}} = Z_{0.975} = 1.96$$

بنابراین فرض H_0 را رد می کنیم اگر $|Z| > 1.96$ باشد یا به عبارتی:

$$|Z| > 1.96 \quad \text{یا} \quad Z < -1.96$$

۲۱-۱۱ حال مقدار آماره آزمون را بدست می آوریم:

ابتدا توجه کنید که مقدار واقعی انحراف معیار طول عمر لامپها را نداریم بنابراین از واریانس یا انحراف معیار $\bar{X}$ برای تخمین واریانس واقعی طول عمر

$$\sigma_{\bar{X}} = \frac{\sigma}{\sqrt{n}} = \frac{120}{\sqrt{100}} = 12$$

لامپها استفاده می کنیم:

به این ترتیب آماره آزمون بصورت زیر بدست می آید:

$$Z = \frac{X' - \mu_0}{\frac{\sigma}{\sqrt{n}}} > \frac{1570 - 1600}{12} = -2.5$$

ملاحظه می کنید که -2.5 خارج از بازه $(-1.96, 1.96)$ قرار دارد بنابراین H_0 را در سطح معنی دار 0.05 رد می کنیم یعنی ادعای مصرف کنندگان در مورد عدم صحت میانگین طول عمر مطرح شده توسط کارخانه سازنده، صحیح می باشد.

ب) حال آزمون فرض را در سطح معنی دار 0.01 انجام می دهیم:

$$\alpha = 0.01 \rightarrow 1 - \frac{\alpha}{2} = 0.995$$

$$\rightarrow Z_{1-\frac{\alpha}{2}} = Z_{0.995} = 2.58$$

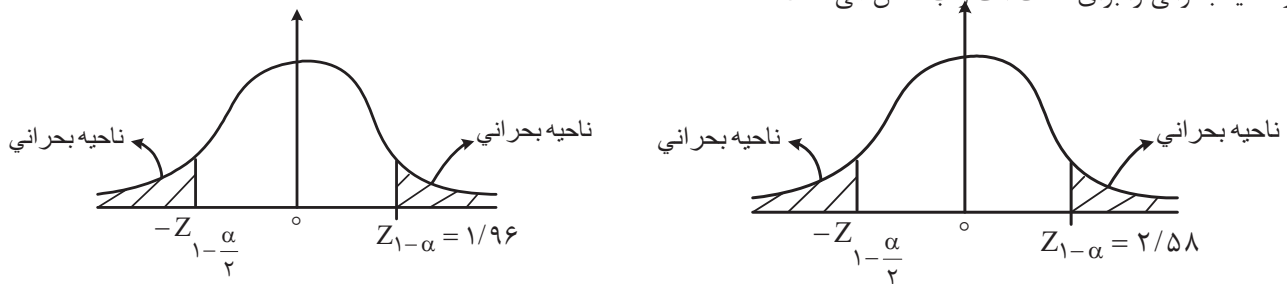
در این حالت اگر آماره آزمون در خارج از بازه $(-2.58, 2.58)$ قرار داشته باشد H_0 را رد می کنیم. مقدار آماره آزمون در این حالت برابر همان مقدار بدست آمده در بند الف مثال می باشد:

$$Z = -2.5 \Rightarrow |Z| > 2.58$$

۲۲-۱۱ از آنجا که مقدار Z برابر 2.5 می باشد و در بازه $(-2.58, 2.58)$ قرار دارد بنابراین در سطح معنی دار 0.01 و فرض H_0 را

می پذیریم.

ملاحظه می‌کنید که با تغییر سطح معنی‌داری و به دنبال آن تغییر ناحیه بحرانی، این احتمال وجود دارد که تصمیم‌گیری در مورد رد یا پذیرش فرض H_0 کاملاً عوض شود. توجه کنید که در حالت دوم ($\alpha = 0.01$) فرض H_0 را می‌پذیریم اما این به معنی رد فرض H_1 نمی‌باشد بلکه در این حالت می‌گوییم نمی‌توان در مورد رد فرض H_1 نظری داد یا به عبارت دیگر هیچ تصمیمی نمی‌گیریم. دو نمودار زیر ناحیه بحرانی را برای حالت الف و ب نشان می‌دهند:



۱۱-۲۳ آزمون فرض برای میانگین نرمال با واریانس نامعلوم

در صورتی که واریانس جامعه نامعلوم باشد نشان دادیم که آماره $T = \frac{\bar{X} - \mu_0}{S / \sqrt{n}}$ دارای توزیع t با $n-1$ درجه آزادی است، بنابراین در این حالت و با

توجه به روشهای ارایه شده برای بدست آوردن ناحیه بحرانی در بخش قبل، می‌توان نشان داد که آماره آزمون و ناحیه بحرانی برای هر حالت بصورت زیر بدست می‌آیند:

$$T = \frac{\bar{X} - \mu_0}{S / \sqrt{n}} > t_{1-\alpha}(n-1) \quad \begin{cases} H_0: \mu = \mu_0 \text{ یا } \mu \leq \mu_0 \\ H_1: \mu > \mu_0 \end{cases} \quad \text{۱- در آزمون فرض } H_0 \text{ رد می‌شود اگر و فقط اگر:}$$

$$T = \frac{\bar{X} - \mu_0}{S / \sqrt{n}} < -t_{1-\alpha}(n-1) \quad \begin{cases} H_0: \mu = \mu_0 \text{ یا } \mu \geq \mu_0 \\ H_1: \mu < \mu_0 \end{cases} \quad \text{۲- در آزمون فرض } H_0 \text{ رد می‌شود اگر و فقط اگر:}$$

$$|T| = \left| \frac{\bar{X} - \mu_0}{S / \sqrt{n}} \right| > t_{1-\frac{\alpha}{2}}(n-1) \quad \begin{cases} H_0: \mu = \mu_0 \\ H_1: \mu \neq \mu_0 \end{cases} \quad \text{۳- در آزمون فرض } H_0 \text{ رد می‌شود اگر و فقط اگر:}$$

۱۱-۲۴ مثال ۷: تعداد نامه‌های رسیده به یک شرکت در طول ۱۰ روز عبارتست از:

۱, ۳, ۲, ۵, ۶, ۴, ۵, ۸, ۹, ۱

اگر توزیع تعداد نامه‌های رسیده از متغیر تصادفی نرمال پیروی کند، آیا در سطح معنی‌دار $\alpha = 0.01$ می‌توان ادعا نمود که بطور متوسط روزانه حداقل ۴ نامه به شرکت پست ارسال می‌شود؟

حل: فرضهای زیر را در نظر می‌گیریم:

$$\begin{cases} H_0: \mu = 4 \\ H_1: \mu > 4 \end{cases}$$

واریانس جامعه نامعلوم است بنابراین برای محاسبه از آماره $T = \frac{\bar{X} - \mu_0}{S / \sqrt{n}}$ استفاده می‌کنیم که در آن $\bar{X}$ و S برابر هستند با:

$$\bar{X} = \frac{1}{10} (1+3+2+5+6+4+5+8+9+1) = 4/4$$

$$S^2 = \frac{1}{n-1} \left(\sum_{i=1}^n X_i^2 - \frac{\bar{X}^2}{n} \right) = \frac{1}{9} \left\{ (1+9+4+25+36+16+25+64+81+1) - \frac{(4/4)^2}{10} \right\}$$

$$= \frac{1}{9} \left(262 - \frac{(4/4)^2}{10} \right) = 28/8 \Rightarrow S = \sqrt{29} = 5/36$$

فرض H_0 در صورتی رد می‌شود که:

$$T > t_{1-\alpha} (n-1)$$

$$\alpha = 0/01 \rightarrow 1-\alpha = 0/99$$

$$t_{1-\alpha} = t_{0/99} (n-1) = t_{0/99} (9) = 2/82$$

مقدار آماره آزمون عبارتست از:

$$T = \frac{\bar{X} - \mu_0}{\frac{S}{\sqrt{n}}} = \frac{4/4 - 4}{\frac{5/36}{\sqrt{10}}} = 0/23$$

از آنجا که مقدار آماره آزمون کمتر از $2/82$ می‌باشد بنابراین فرض H_0 را می‌پذیریم.

۲۵-۱۱. ۹ آزمون فرض برای واریانس یک جامعه نرمال

در صورتی که بخواهیم ادعاهایی را در مورد واریانس یک جامعه نرمال بررسی کنیم می‌دانیم که آماره $\chi^2 = \frac{(n-1) S^2}{\sigma^2}$ دارای توزیع خی دو با $n-1$ درجه آزادی است بنابراین می‌توان حالت‌های زیر را برای هر حالت بدست آورد:

$$\chi^2 = \frac{(n-1) S^2}{\sigma_0^2} > \chi_{1-\alpha}^2 (n-1) \quad H_0: \sigma^2 = \sigma_0^2 \text{ یا } \sigma^2 \leq \sigma_0^2 \quad \text{رد می‌شود اگر و فقط اگر:} \quad 1-\text{در آزمون}$$

$$\chi^2 = \frac{(n-1) S^2}{\sigma_0^2} < \chi_{\alpha}^2 (n-1) \quad H_0: \sigma^2 = \sigma_0^2 \text{ یا } \sigma^2 \geq \sigma_0^2 \quad \text{رد می‌شود اگر و فقط اگر:} \quad 2-\text{در آزمون فرض}$$

$$H_0: \sigma^2 = \sigma_0^2 \quad \text{رد می‌شود اگر و فقط اگر} \quad 3-\text{در آزمون فرض} \quad \begin{cases} H_0: \sigma^2 = \sigma_0^2 \\ H_1: \sigma^2 \neq \sigma_0^2 \end{cases}$$

$$\chi^2 = \frac{(n-1) S^2}{\sigma_0^2} > \chi_{1-\frac{\alpha}{2}}^2 (n-1) \quad \text{یا} \quad \chi^2 = \frac{(n-1) S^2}{\sigma_0^2} < \chi_{\frac{\alpha}{2}}^2 (n-1) :$$

۲۶-۱۱ مثال ۹: یک کارخانه تولید نوشابه ماشین آلات جدیدی برای پرمودن شیشه‌ها خریداری نموده است که ادعا شده است این ماشین‌ها با

انحراف معیار ۱۰ میلی لیتر شیشه‌ها را پر می‌کنند. برای بررسی صحت ادعا یک نمونه تصادفی ۱۰ تایی از شیشه‌های پر شده را انتخاب می‌کنیم و

نتایج $\bar{X} = 255$ و $S^2 = 180$ بدست آمده است. در سطح معنی‌دار $\alpha = 0/05$ آیا ادعای مطرح شده صحیح است؟

حل: فرضها عبارتند از:

$$\begin{cases} H_0 : \sigma^2 = 100 \\ H_1 : \sigma^2 \neq 100 \end{cases}$$

با توجه به آزمونها آماره آزمون و ناحیه بحرانی بصورت زیر خواهند بود:

$$\chi^2_{1-\frac{\alpha}{2}} (n-1) < \chi^2 = \frac{(n-1) S^2}{\sigma_0^2} \quad \text{یا} \quad \chi^2 = \frac{(n-1) S^2}{\sigma_0^2} < \chi^2_{\frac{\alpha}{2}} (n-1)$$

در سطح معنی‌دار ۰/۰۵ داریم:

$$\alpha = 0.05 \quad \rightarrow \quad 1 - \frac{\alpha}{2} = 0.975 \quad \frac{\alpha}{2} = 0.025$$

پس:

$$\chi^2_{1-\frac{\alpha}{2}} (n-1) = \chi^2_{0.975} (9) = 19$$

$$\chi^2_{\frac{\alpha}{2}} (n-1) = \chi^2_{0.025} (9) = 16.919$$

$$\chi^2 = \frac{(n-1) S^2}{\sigma_0^2} = \frac{9(110)}{100} = 9.9$$

از آنجا که آماره $\chi^2 = 9.9$ درون بازه (۱۹ , ۱۶/۹۱۹) قرار دارد بنابراین ادعای H_0 رد نمی‌شود و در نتیجه ادعا مطرح شده در سطح معنی‌دار ۰/۰۵ قابل قبول است.

۹-۱۱-۲۷ آزمون فرض برای تفاضل میانگین‌ها

یک نمونه تصادفی n_1 تایی از جامعه نرمال با میانگین μ_1 و واریانس σ_1^2 انتخاب می‌کنیم، همچنین یک نمونه n_2 تایی از جامعه نرمال دیگری با میانگین μ_2 و واریانس σ_2^2 انتخاب می‌کنیم بطوریکه دو نمونه از یکدیگر مستقل باشند در این حالت برای آزمون تفاضل میانگین‌های دو جامعه از آزمون‌های زیر می‌توانیم استفاده کنیم که در سه حالت اول واریانس دو جامعه معلوم و در سه حالت بعدی واریانس مجهول فرض شده است:

$$Z = \frac{(\bar{X}_1 - \bar{X}_2) - d_0}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} > Z_{1-\alpha} \quad \text{فرض } H_0 : \mu_1 - \mu_2 = d_0 \text{ یا } \mu_1 - \mu_2 \leq d_0 \quad \text{در آزمون فرض } H_1 : \mu_1 - \mu_2 > d_0$$

$$Z = \frac{(\bar{X}_1 - \bar{X}_2) - d_0}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} < -Z_{1-\alpha} \quad \text{فرض } H_0 : \mu_1 - \mu_2 = d_0 \text{ یا } \mu_1 - \mu_2 \geq d_0 \quad \text{در آزمون فرض } H_1 : \mu_1 - \mu_2 < d_0$$

$$|Z| = \left| \frac{(\bar{X}_1 - \bar{X}_2) - d_0}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} \right| > Z_{1-\alpha} \quad \text{فرض } H_0 : \mu_1 - \mu_2 = d_0 \quad \text{در آزمون فرض } H_1 : \mu_1 - \mu_2 \neq d_0$$

۹-۱۱-۲۸ در حالتی که واریانسهای دو جامعه نامعلوم اما برابر باشند ($\sigma_1^2 = \sigma_2^2 = \sigma^2$) روابط آزمونه‌های فرض بصورت زیر خواهند بود:

۱- در آزمون فرض $\begin{cases} H_0 : \mu_1 - \mu_2 = d_0 \text{ یا } \mu_1 - \mu_2 \leq d_0 \\ H_1 : \mu_1 - \mu_2 > d_0 \end{cases}$ فرض H_0 رد می‌شود اگر و فقط اگر:

$$T = \frac{(\bar{X}_1 - \bar{X}_2) - d_0}{S_P \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} > t_{1-\alpha} (n_1 + n_2 - 2)$$

$$\text{که در آن } S_P^2 = \frac{(n_1 - 1) S_1^2 + (n_2 - 1) S_2^2}{n_1 + n_2 - 2} \text{ می‌باشد.}$$

۲- در آزمون فرض $\begin{cases} H_0 : \mu_1 - \mu_2 = d_0 \text{ یا } \mu_1 - \mu_2 \geq d_0 \\ H_1 : \mu_1 - \mu_2 < d_0 \end{cases}$ فرض H_0 رد می‌شود اگر و فقط اگر:

$$T = \frac{(\bar{X}_1 - \bar{X}_2) - d_0}{S_P \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} < t_{1-\alpha} (n_1 + n_2 - 2)$$

۳- در آزمون فرض $\begin{cases} H_0 : \mu_1 - \mu_2 = d_0 \\ H_1 : \mu_1 - \mu_2 \neq d_0 \end{cases}$ فرض H_0 رد می‌شود اگر و فقط اگر:

$$|T| = \left| \frac{(\bar{X}_1 - \bar{X}_2) - d_0}{S_P \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} \right| > t_{1-\frac{\alpha}{2}} (n_1 + n_2 - 2)$$

۲۹-۱۱ مثال ۱۰: یک امتحان از دو کلاس A و B گرفته شده است. کلاس A شامل ۴۰ دانشجو و کلاس B شامل ۵۰ دانشجو می‌باشد. میانگین

نمرات در کلاس A ۷۴ با انحراف معیار ۸ و در کلاس B ۷۸ با انحراف معیار ۷ می‌باشد آیا نمرات دانشجویان این دو کلاس متفاوت معنی‌داری در سطوح ۰/۰۵ و ۰/۰۱ با هم دارند؟

حل: هر یک از دو کلاس را دو جامعه با میانگین‌های μ_1 و μ_2 در نظر می‌گیریم داریم:

$$\bar{X}_1 = 74 \quad \sigma_1 = 8$$

$$\bar{X}_2 = 78 \quad \sigma_2 = 7$$

$$n_1 = 40$$

$$n_2 = 50$$

فرض‌های آزمون عبارتند از:

$H_0 : \mu_1 = \mu_2$ اختلاف تنها ناشی از شانس است.

$H_1 : \mu_1 \neq \mu_2$ بین دو کلاس اختلاف معنی‌داری وجود دارد.

در این حالت آماره آزمون و ناحیه بحرانی عبارتند از:

$$|Z| = \left| \frac{(\bar{X}_1 - \bar{X}_2) - d_0}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} \right| > Z_{1-\alpha}$$

که در این مثال مقدار $d_0 = 0$ می‌باشد.

۳۰-۱۱ الف) در سطح معنی‌دار ۰/۰۵ داریم:

$$\alpha = 0.05 \rightarrow 1 - \frac{\alpha}{2} = 0.975$$

$$Z_{1-\frac{\alpha}{2}} = Z_{0.975} = 1.96$$

در این حالت فرض H_0 رد می‌شود اگر $|Z| > 1.96$ حال داریم:

$$Z = \frac{(\bar{X}_1 - \bar{X}_2) - d_0}{\sqrt{\frac{\sigma_1^2}{n_1} + \frac{\sigma_2^2}{n_2}}} = \frac{74 - 78}{\sqrt{\frac{82}{40} + \frac{72}{50}}} = \frac{-4}{1.606} = -2.49$$

از آنجا که $Z = -2.49$ در بازه $(-1.96, 1.96)$ قرار ندارد بنابراین در سطح معنی‌دار ۰/۰۵ اختلاف معنی‌داری بین دو کلاس A و B وجود دارد و فرض H_0 رد می‌شود. به این ترتیب به احتمال بیشتر کلاس دوم دارای نتایج بهتری می‌باشد.
ب) برای سطح معنی‌دار ۰/۰۱ داریم:

$$\alpha = 0.01 \rightarrow 1 - \frac{\alpha}{2} = 0.995$$

$$Z_{1-\frac{\alpha}{2}} = Z_{0.995} = 2.58$$

در این حالت مقدار آماره آزمون برابر با مقدار بدست آمده از بند الف می‌باشد که برابر -2.49 می‌باشد. و از آنجا که -2.49 در بازه $(-2.58, 2.58)$ قرار دارد بنابراین در سطح معنی‌دار ۰/۰۱ اختلاف معنی‌داری بین دو کلاس وجود ندارد. توجه کنید که در حالت کلی نتایج حاصل از سطح معنی‌دار ۰/۰۵ را ملاک تصمیم‌گیری قرار می‌دهند. زیرا آماردانان نشان داده‌اند که تقریباً بهترین سطح برای ملاک بودن در تصمیم‌گیری‌ها، سطح معنی‌دار ۰/۰۵ می‌باشد.

۳۱-۱۱ مثال ۱۱: در کشاورز A و B در مزرعه‌های خود گندم می‌کارند، کشاورز A از نوعی ضد آفت جدید استفاده می‌کند. میانگین برداشت محصول در هر کیلومتر مربع از ۱۲ کیلومتر مربع از زمین کشاورز A برابر ۱۳۹ کیلو با انحراف معیار ۱۰ کیلو می‌باشد و در زمین کشاورز B برابر ۱۳۱ کیلو با انحراف معیار ۱۱ کیلو می‌باشد. آیا می‌توان ادعا کرد که در سطح معنی‌دار ۰/۰۵ و ۰/۰۱.

حل: فرض آزمون عبارتند از:

$H_0: \mu_1 = \mu_2$: اختلاف در تولید تنها ناشی از شانس است.

$H_1: \mu_1 > \mu_2$: ضد آفت در افزایش تولید موثر بوده است.

μ_1 میانگین تولید محصول کشاورز A و μ_2 میانگین تولید محصول کشاورز B می‌باشد.

آماره آزمون و ناحیه بحرانی در این حالت عبارتند از:

$$T = \frac{(\bar{X}_1 - \bar{X}_2) - d_0}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} > t_{1-\alpha} (n_1 + n_2 - 2)$$

الف) در سطح معنی‌دار ۰/۰۵ داریم:

$$\alpha = 0.05$$

$$1 - \alpha = 0.95 \quad t_{1-\alpha} (n_1 + n_2 - 2) = t_{0.95} (12 + 12 - 2)$$

$$= t_{0.95} (22) = 1.72$$

$$\bar{X}_1 = 139 \quad S_1 = 10 \quad n_1 = 12$$

$$\bar{X}_2 = 131 \quad S_2 = 11 \quad n_2 = 12$$

مقدار S_p برابر است با:

$$S_p = \sqrt{\frac{(n_1 - 1) S_1^2 + (n_2 - 1) S_2^2}{n_1 + n_2 - 2}} = \sqrt{\frac{11(10)^2 + 11(11)^2}{12 + 12 - 2}} = 10.51$$

۳۲-۱۱ حال مقدار آماره آزمون برابر می شود با:

$$T = \frac{\bar{X}_1 - \bar{X}_2}{S_p \sqrt{\frac{1}{n_1} + \frac{1}{n_2}}} = \frac{139 - 131}{10.51 \sqrt{\frac{1}{12} + \frac{1}{12}}} = 1.85$$

با توجه به مقدار آماره آزمون $T = 1.85$ که از $t_{0.95}(22) = 1.72$ بیشتر است می توان نتیجه گرفت که فرض H_0 در سطح معنی دار 0.05 رد می شود و به این ترتیب ضدآفت در تولید محصول بیشتر موثر بوده است.

ب) در سطح معنی دار 0.01 داریم:

$$\alpha = 0.01 \rightarrow 1 - \alpha = 0.99$$

$$t_{1-\alpha}(n_1 + n_2 - 2) = t_{0.99}(22) = 2.51$$

از آنجا که $T = 1.85 \not> 2.51$ بنابراین فرض H_0 رد نمی شود و تفاوت معنی داری در سطح 0.01 در تولید محصول وجود ندارد. اما از آنجا که 0.05 را ملاک تصمیم گیری می گیریم بنابراین در حالت کلی ضد آفت در تولید محصول بیشتر موثر بوده است.

۳۳-۱۱. ۱۰ آزمون فرض برای واریانسهای دو جامعه

اگر دو جامعه نرمال داشته باشیم و n_1 نمونه از جامعه اول با انحراف معیار S_1 و n_2 نمونه با انحراف معیار S_2 از جامعه دوم انتخاب کنیم در این صورت برای انجام آزمونهای در مورد واریانسهای دو جامعه می توانیم از آماره و ناحیه بحرانی زیر استفاده کنیم:

$$F = \frac{S_1^2}{S_2^2} > F_{1-\alpha}(n_1 - 1, n_2 - 1) \quad \text{فرض } H_0 \text{ رد می شود اگر و فقط اگر: } \begin{cases} H_0: \sigma_1^2 = \sigma_2^2 \text{ یا } \sigma_1^2 \leq \sigma_2^2 \\ H_1: \sigma_1^2 > \sigma_2^2 \end{cases} \quad \text{۱- در آزمون فرض}$$

$$F = \frac{S_1^2}{S_2^2} < F_{\alpha}(n_1 - 1, n_2 - 1) \quad \text{فرض } H_0 \text{ رد می شود اگر و فقط اگر: } \begin{cases} H_0: \sigma_1^2 = \sigma_2^2 \text{ یا } \sigma_1^2 \geq \sigma_2^2 \\ H_1: \sigma_1^2 < \sigma_2^2 \end{cases} \quad \text{۲- در آزمون فرض}$$

$$\text{۳- در آزمون فرض } \begin{cases} H_0: \sigma_1^2 = \sigma_2^2 \\ H_1: \sigma_1^2 \neq \sigma_2^2 \end{cases} \quad \text{فرض } H_0 \text{ رد می شود اگر و فقط اگر:}$$

$$F = \frac{S_1^2}{S_2^2} < F_{\frac{\alpha}{2}}(n_1 - 1, n_2 - 1) \text{ یا } F = \frac{S_1^2}{S_2^2} > F_{1-\frac{\alpha}{2}}(n_1 - 1, n_2 - 1)$$

۳۴-۱۱ مثال ۱۲: استادی یک درس را در دو کلاس A و B تدریس می کند. کلاس A تعداد ۱۶ دانشجو و کلاس B تعداد ۲۵ دانشجو دارد. استاد

از هر دو کلاس یک امتحان را می گیرد. میانگین نمرات دانشجویان در هر کلاس تقریباً برابر است اما واریانس کلاس A برابر $86/4$ و واریانس کلاس B برابر 150 می باشد. (نمرات از 100 واحد می باشند) آیا در دو سطح معنی دار 0.05 و 0.01 می توان نتیجه گرفت که واریانس کلاس B از واریانس کلاس A بیشتر می باشد؟

حل: فرضهای آزمون عبارتند از:

$H_0: \sigma_1 = \sigma_2$: اختلاف واریانس ها ناشی از شانس می باشد.

$H_1: \sigma_1 > \sigma_2$ واریانس کلاس B از واریانس کلاس A بیشتر است.

$$\begin{aligned} n_1 &= 16 & S_1^2 &= 86/4 \\ n_2 &= 25 & S_2^2 &= 150 \end{aligned}$$

داریم:

آماره آزمون و ناحیه بحرانی عبارتند از:

$$F = \frac{S_1^2}{S_2^2} < F_{\alpha}(n_1-1, n_2-1)$$

الف) در سطح معنی‌دار ۰/۰۵ داریم:

$$\alpha = 0/05$$

$$\begin{aligned} F_{\alpha}(n_1-1, n_2-1) &= \frac{1}{F_{1-\alpha}(n_2-1, n_1-1)} \\ \Rightarrow F_{0/05}(15, 24) &= \frac{1}{F_{0/95}(24, 15)} = \frac{1}{2/11} = 0/473 \end{aligned}$$

۱۱-۳۵ حال آماره آزمون را بدست می‌آوریم:

$$F = \frac{S_1^2}{S_2^2} = \frac{86/4}{150} = 0/576$$

از آنجا که $f = 0/576 \nless 0/473$ بنابراین نمی‌توانیم فرض H_0 را رد کنیم و در نتیجه می‌توان گفت که واریانس دو کلاس در سطح معنی‌دار ۰/۰۵ تفاوت معنی‌داری با یکدیگر ندارند.

ب) در سطح معنی‌دار ۰/۰۱ داریم:

$$\alpha = 0/01$$

$$\begin{aligned} F_{\alpha}(n_1-1, n_2-1) &= \frac{1}{F_{1-\alpha}(n_2-1, n_1-1)} \\ \Rightarrow F_{0/01}(15, 24) &= \frac{1}{F_{0/99}(24, 15)} = \frac{1}{2/89} = 0/346 \end{aligned}$$

باز هم $f = 0/576 \nless 0/346$ بنابراین در سطح معنی‌دار ۰/۰۱ هم می‌توانیم H_0 را رد کنیم یعنی در این حالت هم تفاوت معنی‌داری میان واریانسهای دو کلاس وجود ندارد.

فصل دوازدهم

S-1

۱- دگرسیون-۱

دگرسیون

در فصلهای قبلی بیشتر تجزیه و تحلیل‌های آماری روی یک صفت از جامعه یا متغیر تصادفی متمرکز بود در این فصل قصد داریم بصورت همزمان دو صفت از جامعه یا دو متغیر تصادفی را مورد بررسی قرار دهیم. این بررسی شامل پیش بینی مقادیر یکی از متغیرها از روی مقادیر متغیر دیگر است که به مساله برگشت یا دگرسیون معروف می‌باشد.

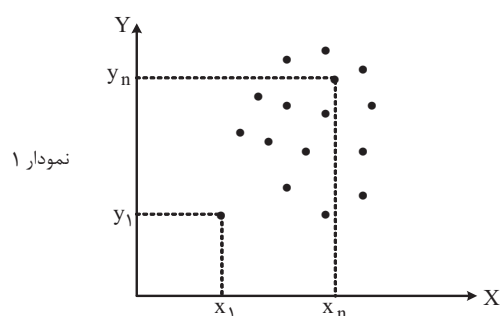
به عنوان مثال فرض کنید بخواهیم تاثیر میزان مصرف شیر را در افزایش قد بدست بیاوریم و یا بخواهیم میزان وزن فرزند را از روی وزن پدرش پیش بینی کنیم. ملاحظه می‌کنید که در این گونه مسایل دو متغیر تصادفی مورد مطالعه به نوعی به یکدیگر وابسته می‌باشند. به عبارت دقیقتر یک متغیر تصادفی مثل X را مستقل و متغیر تصادفی Y را وابسته به آن در نظر می‌گیریم و یا برعکس Y را مستقل و X را وابسته به آن در نظر می‌گیریم. آشکار است که انتخاب هر یک از دو حالت به نوع مساله بستگی دارد.

در مسایل دگرسیون برای یافتن رابطه بین متغیر تصادفی مستقل X و متغیر وابسته Y ابتدا یک نمونه n تایی از متغیر تصادفی X جمع آوری می‌کنیم که نتایج آن بصورت $X_1, X_2, \dots, X_n$ می‌باشند. سپس مقادیر متناظر با هر یک از نمونه‌های بدست آمده (X_i) را که همان مقادیر معادل متغیر تصادفی وابسته Y می‌باشند بدست می‌آوریم. به این ترتیب برای X_i ها مقادیر متناظر $Y_1, Y_2, \dots, Y_n$ بدست می‌آیند. که می‌توانیم نتیجه را بصورت زوج‌های مرتب زیر نشان دهیم.

$$(X_1, Y_1), (X_2, Y_2), (X_3, Y_3), \dots, (X_n, Y_n)$$

۲- دگرسیون-۲ به عنوان مثال مقادیر X_i ها می‌توانند میزان طول قد افراد یک جامعه و مقادیر Y_i ها میزان مصرف شیر هر یک از نمونه‌ها باشند. به این ترتیب زوج مرتب $(180, 2)$ بیانگر این است که در نمونه‌گیری یکی از افراد جامعه دارای صول قد 180cm بوده است و وی روزانه دو لیوان شیر مصرف کرده است.

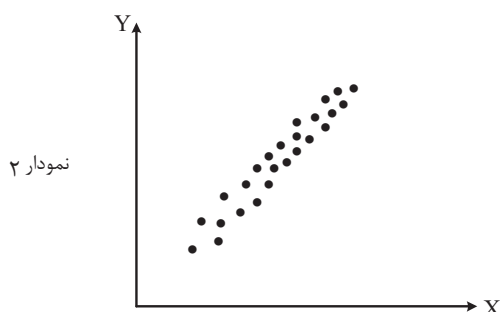
پس از بدست آمدن زوج‌های مرتب (X_i, Y_i) ملاحظه می‌نید که هر زوج مرتب می‌تواند معادل یک نقطه در صفحه باشد. با رسم نقاط مورد نظر در صفحه یک تصویر کلی از رابطه X و Y بدست می‌آوریم. به شکل‌های زیر که برای سه نمونه جداگانه می‌باشند توجه کنید:



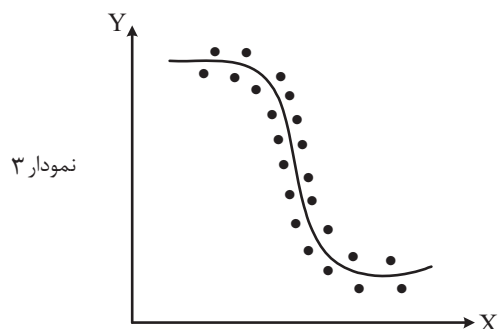
نمودار ۱

در این نمودار ملاحظه می‌کنید که زوج‌های مرتب (X_i, Y_i) بصورت کاملاً پراکنده توزیع شده‌اند و به این ترتیب نتیجه می‌گیریم که رابطه‌ای بین متغیرهای تصادفی X و Y وجود ندارد.

ملاحظه می‌کنید که در این نمونه یک رابطه خطی بین مقادیر X_i و Y_i وجود دارد و می‌توان نوشت $y_i = a x_i + b$.



نمودار ۲



در این نمودار X و Y به یکدیگر وابسته می‌باشند اما این وابستگی از نوع غیر خطی می‌باشد.

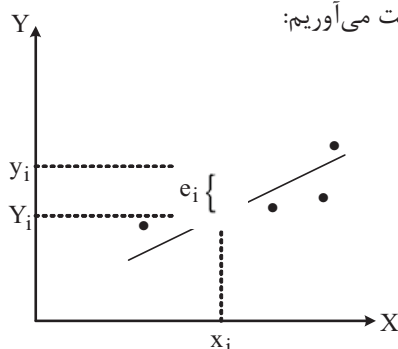
با توجه به دو نمودار ۲ و ۳ این طور به نظر می‌رسد که در حالت کلی دو متغیر تصادفی X و Y اگر از یکدیگر مستقل نباشند یا بصورت خطی و یا بصورت غیر خطی به یکدیگر وابسته می‌باشند.

قبلاً نشان دادیم که برای اندازه‌گیری میزان وابستگی دو متغیر X و Y می‌توان از ضریب همبستگی استفاده نمود. اما در مسایل دگرسیون پیش بینی متغیر Y از روی X و یا بالعکس از اهمیت ویژه‌ای برخوردار است. بنابراین نیازمند روشی هستیم که بتوان در صورت نیاز با ثابت در نظر گرفتن یکی از مقادیر X یا Y مقدار دیگری را بدست بیاوریم. برای این منظور مفهوم بردازش منحنی را مطرح می‌کنیم.

۱- دگرسیون خطی - ۱.۱۲.۳ دگرسیون خطی

اگر بین دو متغیر X و Y یک رابطه خطی وجود داشته باشد می‌توانیم یک خط را طوری رسم کنیم که نقاط (x_i, y_i) کمترین فاصله را با خط مورد نظر داشته باشند. به این عمل بردازش منحنی می‌گویند. معادله خط را بصورت $Y = aX + b$ در نظر می‌گیریم که در آن a و b مقادیر مجهول می‌باشند و در این حالت X متغیر مستقل و Y متغیر وابسته به آن در نظر گرفته می‌شود. مقادیر a و b می‌بایستی طوری محاسبه شوند که مجموع فاصله نقاط (x_i, y_i) از خط $Y = aX + b$ حداقل شود. در این حالت به $Y = aX + b$ معادله دگرسیون Y می‌گویند.

برای حداقل نمودن فاصله نقاط (x_i, y_i) از خط دگرسیون مقدار خطای e_i را مطابق نمودار زیر بدست می‌آوریم:



با توجه به نمودار Y_i مقدار پیش بینی شده توسط خط دگرسیون می‌باشد که با مقدار واقعی y_i به اندازه $e_i = |y_i - Y_i|$ فاصله دارد که این فاصله همان خطای پیش بینی می‌باشد.

برای بدست آوردن بهترین نتیجه، مجموع مربعات خطا را حداقل می‌کنیم که عبارتست از:

$$SSE = E = \sum_{i=1}^n (y_i - Y_i)^2$$

که در آن:

$$Y = aX + b \Rightarrow Y_i = aX_i + b \Rightarrow E = \sum_{i=1}^n (aX_i + b - Y_i)^2$$

۲- دگرسیون خطی - ۴ برای مینیمم نمودن E مشتقات پاره‌ای آنرا نسبت به a و b بدست آورده و برابر صفر قرار می‌دهیم:

$$\frac{\partial E}{\partial a} = \sum_{i=1}^n \gamma x_i (a x_i + b - y_i) = 0$$

$$\frac{\partial E}{\partial b} = \sum_{i=1}^n \gamma (a x_i + b - y_i) = 0$$

به این ترتیب دستگاه معادلات زیر بدست می‌آید:

$$\begin{cases} \sum_{i=1}^n \gamma x_i (a x_i + b - y_i) = 0 \\ \sum_{i=1}^n \gamma (a x_i + b - y_i) = 0 \end{cases}$$

با حل این دستگاه مقادیر مجهول a و b بدست می‌آیند. که عبارتند از:

$$a = \frac{S_{xy}}{S_{xx}}, \quad b = \bar{y} - a \bar{X}$$

که در آن:

$$S_{xy} = \sum_{i=1}^n (x_i - \bar{x})(y_i - \bar{y}) = \sum_{i=1}^n x_i y_i - n \bar{x} \bar{y}$$

$$S_{xx} = \sum_{i=1}^n (x_i - \bar{x})^2 = \sum_{i=1}^n x_i^2 - n \bar{x}^2$$

به معادله $y = a x + b$ خط دگرسیون Y روی x می‌گویند. یعنی Y از روی x بدست آمده است.

۱-مثال ۵-مثال ۱: در یک بررسی میزان قد پدران و فرزندان آنها بصورت جدول زیر بدست آمده است:

قد پدران X	۱۶۵	۱۶۰	۱۷۰	۱۶۳	۱۷۳	۱۵۷	۱۷۸	۱۶۸	۱۷۳	۱۷۰	۱۷۵	۱۸۰
قد فرزندان Y	۱۷۳	۱۶۸	۱۷۳	۱۶۵	۱۷۵	۱۶۸	۱۷۳	۱۶۵	۱۸۰	۱۷۰	۱۷۳	۱۷۸

مطلوبست:

الف) برداش خطی داده‌ها

ب) رسم نمودار داده‌ها و خط دگرسیون.

ج) با توجه به معادله دگرسیون پیش بینی کنید که اگر قد پدری ۱۸۵cm باشد قد فرزند او چقدر خواهد بود؟

حل: برای بدست آوردن a و b در معادله $y = a x + b$ جدول زیر را تشکیل می‌دهیم:

x	y	x^2	xy	y^2
۱۶۵	۱۷۳	۲۷۲۲۵	۲۸۵۴۵	۲۹۹۲۹
۱۶۰	۱۶۸	۲۵۶۰۰	۲۶۸۸۰	۲۸۲۲۴
۱۷۰	۱۷۳	۲۸۹۰۰	۲۹۴۱۰	۲۹۹۲۹
۱۶۳	۱۶۵	۲۶۵۶۹	۲۶۸۹۵	۲۷۲۲۵
۱۷۳	۱۷۵	۲۹۹۲۹	۳۰۲۷۵	۳۰۶۲۵
۱۵۸	۱۶۸	۲۴۹۶۴	۲۶۵۴۴	۲۸۲۲۴
۱۷۸	۱۷۳	۳۱۶۸۴	۳۰۷۹۴	۲۹۹۲۹
۱۶۸	۱۶۵	۲۸۲۲۴	۲۷۷۲۰	۲۷۲۲۵
۱۷۳	۱۸۰	۲۹۹۲۹	۳۱۱۴۰	۳۲۴۰۰
۱۷۰	۱۷۰	۲۸۹۰۰	۲۸۹۰۰	۲۸۹۰۰

۲- مثال ۱-۶

$$\bar{X} = \frac{1}{n} \sum x_i = \frac{1}{12} 2033 = 169/41$$

$$\bar{Y} = \frac{1}{n} \sum y_i = \frac{1}{12} 2061 = 171/75$$

$$S_{xy} = \sum_{i=1}^n x_i y_i - n \bar{x} \bar{y} = 349418 - 12(169/41)(171/75) = 250/25$$

$$S_{xx} = \sum x_i^2 - n \bar{x}^2 = 344949 - 12(169/41)^2 = 524/91$$

$$a = \frac{S_{xy}}{S_{xx}} = \frac{250/25}{524/91} = 0/477$$

$$b = \bar{Y} - a \bar{X} = 171/75 - 0/477(169/41) = 90/9$$

بنابراین معادله خط دگرسیون عبارتست از:

$$y = 0/477 x + 90/9$$

ب) نمودار داده‌ها و خط دگرسیون عبارتست از:

ج) با توجه به خط دگرسیون می‌توان قد فرزند یک پدر با قد ۱۸۵ را به فرم زیر بدست آورد.

$$y = 0/477 \times 185 + 90/9 = 179/145$$

۳- S ۱۲. ۲ ضریب همبستگی و دگرسیون

۷-۱

قبلاً ضریب همبستگی دو متغیر X و Y را بصورت زیر تعریف نمودیم:

$$\rho_{XY} = \frac{\sigma_{XY}}{\sigma_X \sigma_Y}$$

می‌توان با ساده نمودن معادله دگرسیون مقدار ضریب همبستگی را وارد معادله نمود:

$$y = a x + b \quad a = \frac{S_{xy}}{S_{xx}} = \frac{\sigma_{xy}}{\sigma_x \sigma_y} \left(\frac{\sigma_y}{\sigma_x} \right) = \rho_{xy} \frac{\sigma_y}{\sigma_x}$$

توجه کنید که در اینجا در محاسبه ρ_{xy} و σ_x و σ_y از مقادیر نمونه‌های مشاهده شده استفاده می‌شود.

به این ترتیب معادله خط دگرسیون بصورت زیر بدست می‌آید:

$$y - \bar{y} = \rho_{xy} \frac{\sigma_y}{\sigma_x} (x - \bar{x})$$

توجه کنید که همواره برای داده‌های مشاهده شده (x_i, y_i) دو معادله دگرسیون وجود دارد یک معادله بر حسب y نسبت به x می‌باشد و معادله

دیگر بر حسب x نسبت به y می‌باشد معادله خط دگرسیون x روی y را می‌توان بصورت زیر بدست آورد:

$$x = \alpha y + \beta$$

$$\alpha = \frac{S_{xy}}{S_{yy}} \quad \beta = \bar{x} - \alpha \bar{y}$$

$$S_{yy} = \sum_{i=1}^n (y_i - \bar{y})^2 = \sum_{i=1}^n y_i^2 - n \bar{y}^2$$

۲-۸ با نوشتن معادله دگرسیون x روی y و استفاده از ضریب همبستگی بدست می‌آوریم:

$$x - \bar{x} = \rho_{xy} \frac{\sigma_x}{\sigma_y} (y - \bar{y})$$

بنابراین در حالت کلی دو خط دگرسیون y روی x و x روی y خواهیم داشت که عبارتند از:

$$y - \bar{y} = \rho_{xy} \frac{\sigma_y}{\sigma_x} (x - \bar{x})$$

$$x - \bar{x} = \rho_{xy} \frac{\sigma_x}{\sigma_y} (y - \bar{y})$$

با برابر قرار دادن $x = y$ در معادلات فوق محل تلاقی دو خط دگرسیون نقطه $\begin{cases} x = \bar{x} \\ y = \bar{y} \end{cases}$ بدست می‌آید. همچنین توجه کنید که:

$$a = \rho_{xy} \frac{\sigma_y}{\sigma_x}$$

$$\Rightarrow a \alpha = \rho_{xy} \frac{\sigma_y}{\sigma_x} \rho_{xy} \frac{\sigma_x}{\sigma_y} = \rho_{xy}^2$$

$$\alpha = \rho_{xy} \frac{\sigma_x}{\sigma_y}$$

بنابراین $a \alpha = \rho_{xy}^2$ و می‌بایستی مقداری در بازه $[-1, 1]$ داشته باشد.

۲-۹ مثال ۲: در مثال ۱ مطلوبست:

الف) محاسبه ضریب همبستگی نمونه با توجه به مقادیر محاسبه شده a و b .

ب) معادله دگرسیون x روی y از روی ضریب همبستگی.

ج) رسم نمودار داده‌ها و دو خط دگرسیون y روی x و x روی y .

د) پیش بینی قد پدر اگر قد فرزند وی ۱۷۹ cm باشد.

حل: الف) ضریب همبستگی، با توجه به روابط ارائه شده عبارتست از:

$$a = \rho_{xy} \frac{\sigma_y}{\sigma_x} = 0.477$$

$$S_{yy} = \sum y_y^2 - n \bar{y}^2 = 354223 - 12(171/75)^2 = 246/25$$

$$\sigma_y = \sqrt{S_{yy}} = \sqrt{246/25} = 15/69$$

$$\sigma_x = \sqrt{S_{xx}} = \sqrt{524/91} = 22/91$$

$$\Rightarrow a = 0.477 = \rho_{xy} \frac{15/69}{22/91} \Rightarrow \rho_{xy} = 0.696$$

ب) معادله دگرسیون X روی Y با توجه به ρ_{xy} برابر است با:

$$x = \alpha y + \beta = \rho_{xy} \frac{\sigma_x}{\sigma_y} (y - \bar{y}) + \bar{x} = 0.696 \left(\frac{22/91}{15/69} \right) (y - 171/75) + 169/41$$

$$\Rightarrow x = 1/0.16 y - 5/12$$

ج) نمودار داده‌ها و دو خط دگرسیون و روی X و X روی Y عبارتست از:

د) با استفاده از خط دگرسیون X روی Y مقدار X را برای $y = 179$ پیش‌بینی می‌کنیم:

$$x = 1/0.16 y - 5/12 = 1/0.16 (179) - 5/12 = 176/744$$

۱۲.۳ دگرسیون منحنی‌های چند جمله‌ای

با تعمیم روشی که برای برازش داده‌ها بصورت خطی ارائه شد به سادگی می‌توان به داده‌ها یک منحنی چند جمله‌ای بفرم

$$y = a_n x^n + a_{n-1} x^{n-1} + \dots + a_1 x + a_0$$

برازش داد. این کار را برای یک منحنی درجه دوم نشان می‌دهیم:

$$Y = a x^2 + b x + c$$

$$E = \sum_{i=1}^n (y_i - Y_i)^2 = \sum_{i=1}^n (a x_i^2 + b x_i + c - y_i)^2$$

حال مجدداً با مشتق گرفتن نسبت به مقادیر مجهول a و b و c و برابر صفر قرار دادن معادلات یک دستگاه بدست می‌آوریم که با حل آن a و b و c بدست می‌آیند:

$$\frac{\partial E}{\partial a} = \sum_{i=1}^n x_i (a x_i + b x_i + c - y_i) = 0$$

$$\frac{\partial E}{\partial b} = \sum_{i=1}^n x_i (a x_i + b x_i + c - y_i) = 0$$

$$\frac{\partial E}{\partial c} = \sum_{i=1}^n (a x_i + b x_i + c - y_i) = 0$$

با داشتن مقادیر نمونه‌های (x_i, y_i) به سادگی می‌توان دستگاه فوق را حل نموده و مقادیر مجهول a و b و c را بدست آورد.

۱۱ – estenbade amari

استنباط آماری بر روی ضرایب دگرسیونی

در بخش قبل بر اسا نمونه $(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)$ مدل ساده دگرسیونی را معرفی نمودیم و به کمک روش کمترین مربعات

$$y_i = a + b x_i + E_i \quad \text{و} \quad b \text{ برآوردگرهای } \hat{b} = \frac{S_{xy}}{S_{xx}}, \quad \hat{a} = \bar{y} - \hat{b} \bar{x} \text{ را معرفی نمودیم.}$$

در این بخش می‌خواهیم مدل احتمالی دگرسیون را معرفی کرده و بر اساس آن برخی از فرمهای آماری را بر روی a و b آزمون کنیم. فرض کنید در مدل ساده دگرسیونی خطا یعنی E_i متغیر تصادفی با توزیع نرمال به فرم زیر باشد:

$$E_i \sim N(0, \sigma^2)$$

همچنین متغیرهای تصادفی E_i به ازای $i = 1, 2, \dots, n$ را مستقل از هم در نظر می‌گیریم در این صورت چون Y_i یک ترکیب خطی از متغیر تصادفی E_i می‌باشد بنابراین:

$$Y_i \sim N(a + b x_i, \sigma^2)$$

۱۲ – estenbade amari

محاسبه برآوردگرهای a و b و σ^2

به ازای هر مقدار ثابت X_i تابع چگالی متغیر تصادفی Y_i برابر سات با:

$$f_{Y_i|X_i}(y_i|x_i) = \frac{1}{\sigma\sqrt{2\pi}} e^{-\frac{1}{2\sigma^2} (Y_i - (a + b x_i))^2} \quad -\infty < y_i < \infty$$

می‌توان به کمک روش ماکزیمم درستنمایی که روش دیگری است ۶ برآوردگرهای a و b و σ بذای بدست آوردن برآوردگرهای ماکزیمم درستنمایی پارامترهای a و b و σ به کمک نمونه تصادفی $(X_1, Y_1), (X_2, Y_2), \dots, (X_n, Y_n)$ ابتدا از تابع درستنمایی (یا لگاریتم آن که که ساده‌تر است) نسبت به a و b و σ مشتق می‌گیریم.

عبارت‌های حاصل را برابر صفر قرار می‌دهیم، و سپس دستگاه معادلات حاصل را حل می‌کنیم.

بنابراین با مشتق گیری جزئی از:

$$\ln(L) = -n \ln(G) - \frac{n}{2} \ln(2\pi) - \frac{1}{2\sigma^2} \sum_{i=1}^n [Y_i - (a + b x_i)]^2$$

نسبت به a و b و σ و برابر گذاشتن عبارت‌های حاصل ۶ به دست می‌آوریم:

$$\frac{\partial \text{Ln}(L)}{\partial a} = \frac{1}{\sigma^2} + \sum_{i=1}^n [Y_i - (a + b x_i)] = 0$$

$$\frac{\partial \text{Ln}(L)}{\partial b} = \frac{1}{\sigma^2} + \sum_{i=1}^n X_i [Y_i - (a + b x_i)] = 0$$

$$\frac{\partial \text{Ln}(L)}{\partial \sigma} = \frac{-n}{\sigma} + \frac{1}{\sigma^3} + \sum_{i=1}^n [Y_i - (a + b x_i)]^2 = 0$$

از دو معادله اول $\hat{a}$ و $\hat{b}$ مشابه با برآوردهای به روش کمترین مربعات به صورت زیر به دست می‌آید.

$$\hat{b} = \frac{S_{xy}}{S_{xx}}$$

$$\hat{a} = \bar{Y} - \hat{b} \bar{X}$$

با قرار دادن $\hat{a}$ و $\hat{b}$ در معادله سوم برآوردهای $\hat{\sigma}^2$ برابر می‌شود:

$$\hat{\sigma}^2 = \frac{S_{yy} - \hat{b} S_{xy}}{n} = \frac{\text{SSE}}{n}$$

اگر در برآوردهای ماکزیمم درست‌نمایی $\hat{\sigma}^2$ ، n را به $n-2$ تبدیل کنیم آنگاه برآوردهای $S^2 = \frac{\text{SSE}}{n-2}$ نااریب خواهد شد یعنی در این حالت:

$$E[S^2] = E\left[\frac{\text{SSE}}{n-2}\right] = \sigma^2 \text{ می‌شود.}$$

۱۳ مثال

مثال ۱۳: فرض کنیم یک کمپانی می‌خواهد تأثیر تبلیغات را در فروش کالاهای تولید شده بررسی کند بدین منظور داده‌های زیر را بعد از ۱۰ ماه بدست می‌آورد. این داده‌ها در جدول زیر مرتب شده است.

به کمک داده‌های بدست آمده برآورد a و b و σ^2 را بدست آورید.

حل:

$$S_{xx} = \sum_{i=1}^n x_i^2 - \frac{(\sum_{i=1}^n x_i)^2}{n} = 9/28 - \frac{(9/4)^2}{10} = 0/444$$

$$S_{xy} = \sum_{i=1}^n x_i y_i - \frac{(\sum_{i=1}^n x_i)(\sum_{i=1}^n y_i)}{n} = 924/8 - \frac{(9/4)(959)}{10} = 23/34$$

$$\bar{Y} = \frac{\sum_{i=1}^n Y_i}{n} = \frac{959}{10} = 95/9, \quad \bar{X} = \frac{\sum_{i=1}^n X_i}{n} = \frac{9/4}{10} = 0/94$$

پس:

$$\hat{b} = \frac{S_{xy}}{S_{xx}} = \frac{23/34}{0/444} = 52/5676 \approx 52/57$$

$$\hat{a} = \bar{Y} - \hat{b}\bar{X} = 95/9 - (52/5676)(0/94) \approx 46/49$$

بنابراین معادله خط دگرسیون به صورت: $y = 46/49 + 52/57 x$ می‌شود.

۱۴ Mesal

برای محاسبه برآوردگر نااریب σ^2 یعنی S^2 به صورت زیر عمل می‌کنیم:

$$S_{yy} = \sum_{i=1}^n (y_i - \bar{y})^2 = \sum_{i=1}^n y_i^2 - \frac{(\sum_{i=1}^n y_i)^2}{n} = 93/569 - \frac{(959)^2}{10} = 1600/9$$

با جایگذاری S_{yy} و $\hat{b}$ و S_{xy} در فرمول SSE خواهیم داشت:

$$SSE = S_{yy} - \hat{b} S_{xy} = 1600/9 - (52/5676)(23/34) = 3700$$

پس:

$$S^2 = \frac{SSE}{n-2} = \frac{373/97}{8} = 46/750$$

حال از برآوردگرهای ماکزیمم درست‌نمایی ضرایب دگرسیون ساده در آزمون فرضیهایی درباره a و b و در ساختن فاصله‌های اطمینان برای این پارامترها استفاده می‌کنیم. بدین منظور نتایج زیر را بدون اثبات می‌پذیریم:

$$\hat{a} \sim N\left(a, \frac{\sigma^2 \sum_{i=1}^n X_i^2}{n S_{xx}}\right) \quad (1)$$

$$\hat{b} \sim N\left(b, \frac{\sigma^2}{S_{xx}}\right) \quad (2)$$

$$\frac{SSE}{\sigma^2} = \frac{(n-2) S^2}{\sigma^2} \sim X_{(n-2)}^2 \quad \text{اگر } S^2 = \frac{SSE}{n-2} \quad (3)$$

(۴) $S^2, \hat{a}$ مستقل از هم هستند.

(۵) $S^2, \hat{b}$ مستقل از هم هستند.

۱۵ – estenbate amari

استبای آماری بر روی b

رابطه خطی بین X و Y اولین موضوعی است که مورد بررسی قرار می‌گیرد. بنابراین علاقمندیم بدانیم که آیا داده‌ها به اندازه کافی اطلاعات در رابطه با ارتباط خطی بودن Y با X می‌دهد یا نه؟ آزمون فرضهای آماری که بر روی پارامتر b ساخته می‌شود ارتباط بین Y و X را روشن می‌کند بعنوان مثال اگر Y با افزایش یا کاهش X تغییر می‌کند می‌توان فرض $b = 0$ را در مقابل فرض $b \neq 0$ آزمون کرد. با توجه به اینکه برآوردگر $\hat{b}$ دارای توزیع

$$\hat{b} \sim N(b, \frac{\sigma^2}{S_{xx}}) \quad T = \frac{\hat{b} - b_0}{\frac{S}{\sqrt{S_{xx}}}} \sim t(n-2)$$

برای آزمون فرضهای

$$\begin{array}{lll} H_0: b = b_0 & H_0: b = b_0 & H_0: b = b_0 \\ (1) & (2) & (3) \\ H_1: b \neq b_0 & H_1: b > b_0 & H_1: b < b_0 \end{array}$$

استفاده نمود.

نواحی بحرانی برای آزمون فرضهای بالا در سطح معنی‌دار α را می‌توان در جدول زیر خلاصه نمود.

هم‌چنین می‌توان با استفاده از تست آماری T با ضریب اطمینان $1 - \alpha$ فاصله اطمینان زیر را برای پارامتر b معرفی نمود.

$$P\left(\hat{b} - t_{\frac{1-\alpha}{2}} \frac{S}{\sqrt{S_{xx}}} < b < \hat{b} + t_{\frac{1-\alpha}{2}} \frac{S}{\sqrt{S_{xx}}}\right) = 1 - \alpha$$

۱۶- مثال ۱-۴

مثال ۴: با توجه به داده‌های مثال ۱ تعیین کنید که آیا پارامتر b اختلاف معنی‌داری از مقدار ۰ دارد یا نه (بوسیله استفاده از یک مدل خطی بین فروش ماهانه و مقدار تبلیغات)

حل: می‌خواهیم آزمون فرض زیر را انجام دهیم:

$$\begin{array}{l} H_0: b = b_0 \\ H_1: b \neq b_0 \end{array}$$

از تست آماری:

$$T = \frac{\hat{b} - b_0}{\frac{S}{\sqrt{S_{xx}}}} \sim t_{\frac{1-\alpha}{2}}(n-2) = t_{\frac{1-\alpha}{2}}(18-2) = t_{\frac{1-\alpha}{2}}(16)$$

استفاده می‌کنیم. با استفاده از جدول توزیع T به ازای $\alpha = 0.05$ $t_{0.025/16} = 2.1098$ بنابراین چون:

$$T = \frac{\hat{b}}{S} = \frac{52/57}{6/84} = 5/12 > 2/30.6$$

باشد پس فرض H_0 را رد می‌کنیم. بنابراین بر اساس این مشاهدات می‌توان گفت که هزینه‌های مربوط به تبلیغات تأثیر در پیش بینی فروش ماهانه کالا دارد. با توجه به اطلاعات داده شده می‌توان یک فاصله اطمینان با ضریب اطمینان $1-\alpha = 0.95$ نیز برای یک پارامتر b بدست آورد.

جایگذاری مقادیر $\hat{b}$ و S و S_{xx} و $t_{(1-\frac{\alpha}{2})}^{(n-2)}$ در $\hat{b} \pm t_{(1-\frac{\alpha}{2})}^{(n-2)} \times \frac{S}{\sqrt{S_{xx}}}$ اطمینان بدست می‌آید.

$$52/57 \pm 2/30.6 \frac{6/84}{\sqrt{0.444}} \Rightarrow 52/57 \pm 23/67 \Rightarrow [28, 90, 76]$$

این فاصله نشان می‌دهد که با احتمال 0.95 مقدار پارامتر b که بوسیله برآوردگر $\hat{b}$ برآورد می‌شود را دارا می‌باشد.

۱۷-مثال ۲-۴

بطور مشابه می‌توان برای پارامتر a فاصله اطمینان و آزمون فرض ساخت. با توجه به اینکه برآوردگر $\hat{a}$ دارای توزیع $\hat{a} \sim N(a, \frac{\sigma^2 \sum x_i^2}{n S_{xx}})$

می‌باشد می‌توان از تست آماری $T = \frac{\hat{a} - a}{S \sqrt{\frac{\sum x_i^2}{n S_{xx}}}} \sim t_{(n-2)}$ استفاده نمود و خواص بحرانی را برای فرض‌های آماری مختلف بر روی پارامتر b

در سطح معنی‌دار α در جدول زیر ساخت. همچنین یک فاصله اطمینان $(1-\alpha) 100\%$ برای a بصورت:

$$P\left(\hat{a} - t_{(1-\frac{\alpha}{2})}^{(n-2)} S \sqrt{\frac{\sum x_i^2}{n S_{xx}}} < a < \hat{a} + t_{(1-\frac{\alpha}{2})}^{(n-2)} S \sqrt{\frac{\sum x_i^2}{n S_{xx}}}\right) = 1-\alpha = 0$$

برآورد $E[Y|X]$ مقدار میانگین Y به شرط مقدار داده شده X .

برآورد میانگین Y به شرط مقدار داده شده X یکی از مسائل مهم است که می‌بایست مورد بررسی قرار گیرد. بعنوان مثال ممکن است مسئول حفاظت جان کارگران کارخانه علاقمند باشد که متوسط حوادثی که برای کارگران اتفاق می‌افتد، به ازای ساعات خاصی از آموزش که به کارگران داده می‌شود را برآورد کند.

فرض کنید که X و Y یک رابطه خطی بر اساس مدل احتمالی دگرسیونی داشته باشد بطوریکه:

$$E[Y|X] = a + bx$$

دیدیم که تابع چگالی شرطی متغیر تصادفی Y به ازای مقدار داده شده X برابر است با:

$$f_{Y|X}(y|x) = \frac{1}{\sigma \sqrt{2\pi}} e^{-\frac{1}{2\sigma^2}(y-(a+bx))^2} \quad -\infty < y < \infty$$

۱۸-مثال ۳-۴

حال می‌خواهیم با فرض خطی بودن رابطه بین X و Y فاصله اطمینان و آزمون فرض برای پارامتر $E[Y|X] = a + bx$ به ازای مقدار داده شده

x بدست آوریم. همانطور که قبلاً اشاره شد $a + bx_p$ میانگین شرطی Y به شرط $X = x_p$ است و از $\hat{Y} = \hat{a} + \hat{b}x_p$ به عنوان یک برآورد برای

$a + bx_p$ یاد کردیم. حال می‌خواهیم فرض $H_0: E[Y|X=x_p] = E_0$ را در مقابل فرض $H_1: E[Y|X=x_p] \neq E_0$ آزمون می‌کنیم

می‌توانیم با توجه به تابع توزیع $\hat{Y}$ تست آماری زیر را برای آزمون فرض فوق معرفی نماییم.

$$T = \frac{\hat{Y} - E_0}{S_{\hat{Y}}} = \frac{\hat{Y} - E_0}{\sqrt{S^2 \left[\frac{1}{n} + \frac{(x_p - \bar{x})^2}{S_{xx}} \right]}} \sim t(n-2)$$

متغیر تصادفی T دارای توزیع t با $(n-2)$ درجه آزادی می‌باشد. بنابراین در سطح معنی‌دار α فرض H_0 زمانی رد می‌شود که $|T| > t_{\frac{\alpha}{2}}^{(n-2)}$

باشد. بطور مشابه می‌توان فرض H_0 داده شده را در مقابل فرض‌های $H_1: E[Y|X=x_p] > E_0$ یا $H_1: E[Y|X=x_p] < E_0$ آزمون نمود که در این صورت فرض H_0 به ترتیب زمانی رد خواهد شد که داسته باشیم:

$$T < -t_{\frac{\alpha}{2}}^{(n-2)} \quad , \quad T > t_{\frac{\alpha}{2}}^{(n-2)}$$

از تست آماری فرض شده می‌توان برای پیدا کردن فاصله اطمینان برای $E[Y|X=x_p]$ نیز استفاده نمود.

$$P\left(\hat{Y} - t_{\frac{\alpha}{2}}^{(n-2)} \sqrt{S^2 \left[\frac{1}{n} + \frac{(x_p - \bar{x})^2}{S_{xx}} \right]} < a + b x_p < \hat{Y} + t_{\frac{\alpha}{2}}^{(n-2)} \sqrt{S^2 \left[\frac{1}{n} + \frac{(x_p - \bar{x})^2}{S_{xx}} \right]}\right)$$

۱۹-مثال

مثال: از داده‌های مثال ۲ استفاده کرده و یک فاصله اطمینان با ضریب اطمینان ۹۵٪ برای $X=1/0$ پیدا کنید.

حل:

ابتدا $\hat{Y}$ را در نقطه $X_p=1/0$ تخمین می‌زنیم:

$$\hat{Y} = \hat{a} + \hat{b} x_p$$

در مثالهای قبل $\hat{a}=46/49$ و $\hat{b}=52/57$ را بدست آوردیم.

$$\hat{Y} = 46/49 + (52/57)(1/0) = 99/06$$

پس:

فرمول مناسب برای فاصله اطمینان برابر بود با:

$$\hat{Y} \pm t_{\frac{\alpha}{2}}^{(n-2)} \sqrt{S^2 \left[\frac{1}{n} + \frac{(x_p - \bar{x})^2}{S_{xx}} \right]}$$

با جایگذاری مقادیر مناسب در عبارت فوق که قبلاض محاسبه شده است خواهیم داشت:

$$99/06 \pm (2/306) \sqrt{46/75 \left[\frac{(1/0 - 0/94)}{0/44} \right]}$$

بعد از محاسبات انجام شده حاصل می‌شود:

$$99/06 \pm 5/18 \Rightarrow [93/88, 10/24]$$

می‌توان از محاسبات فوق نتیجه گرفت که برای هر واحد هزینه تبلیغات ($\$ 10/000$) متوسط فروش ماهانه با احتمال ۹۵٪ در فاصله $\$ 938800$

و $\$ 1042400$ خواهد بود. در شکل زیر برای $E[Y|X]$ رسم شده است.

ضریب همبستگی

در بسیاری از اوقات نیاز به شاخصی داریم که چگونگی ارتباط بین دو متغیر X و Y را اندازه بگیرد. این شاخص ضریب همبستگی خطی بین دو متغیر X و Y نامیده می‌شود.

ضریب همبستگی خطی نمونه‌ای برآورد نامناسب برای ضریب همبستگی خطی تعریف شده در فصل‌های قبل است این معیار میزان قوت ارتباط خطی بین متغیرهای X و Y را اندازه می‌گیرد و اولین بار دانشمند معروفی انگلیسی به نام کارل پیرسون آنرا معرفی نمود از این جهت به آن ضریب همبستگی خطی پیرسون نیز می‌گوییم و به فرم زیر تعریف می‌شود:

$$r = \frac{S_{xy}}{\sqrt{S_{xx} S_{yy}}}$$