

فصل ۵

آزمون‌های نیکویی برازش و مستقل بودن داده‌ها

۱.۵ مقدمه

برای اینکه بدانیم مجموعه‌ای از داده‌ها دارای توزیع خاصی هستند یا نه می‌توانیم از آزمون نیکویی برازش استفاده کنیم. برای توزیع‌های گسسته می‌توان این آزمون را با استفاده از آماره پیرسن X^2 ، آماره نسبت درستنمایی G^2 یا آماره کولموگروف^۱ - اسمیرنوف^۲ برای توزیع‌های پیوسته انجام داد. در این فصل یک آماره جدید مبتنی بر فاصله آنتروپی معرفی می‌کنیم که برای آزمون نیکویی برازش نمونه‌های با اندازه بزرگ مناسب است. نتایج حاصل از آماره آزمون جدید را با نتایج حاصل از برخی آماره‌های دیگر مقایسه می‌کنیم. همچنین یک آماره جدید برای بررسی مستقل بودن داده‌های آماری معرفی خواهیم کرد. آنتروپی نسبی برای اولین بار در سال ۱۹۵۱ توسط کولبک^۳ و لیبلر^۴ معرفی شد که تحت عناوین مختلفی از جمله فاصله کولبک-لیبلر یا واگرایی داده‌ها شناخته شده است. در ادامه به معرفی آنتروپی نسبی و اطلاعات متقابل می‌پردازیم. آنتروپی نسبی اندازه‌گیری فاصله بین دو توزیع می‌باشد که با استفاده از امید ریاضی لگاریتم نسبت درستنمایی به دست می‌آید. به عبارت دیگر، آنتروپی نسبی $D(q||p)$ معیاری برای ناکارآمدی توزیع p در مقابل توزیع واقعی q می‌باشد.

تعریف ۱.۱.۵. آنتروپی نسبی یا فاصله کولبک-لیبلر بین توابع چگالی احتمال $p(x)$ و $q(x)$ به صورت زیر می‌باشد

$$D(q||p) = \sum_{x \in S(x)} q(x) \log_2 \left(\frac{q(x)}{p(x)} \right) = E \left[\log_2 \left(\frac{q(x)}{p(x)} \right) \right], \quad (1.5)$$

که در آن $S(x)$ تکیه‌گاه متغیر تصادفی X است.

با استفاده از خاصیت پیوسته بودن توابع چگالی احتمال داریم

$$\begin{aligned} q(x) = 0 &\implies 0 \log \left(\frac{0}{p} \right) = 0, \\ p(x) = 0 &\implies q \log \left(\frac{q}{0} \right) = \infty. \end{aligned}$$

^۱Kolmogorov

^۲Smirnov

^۳Kullback

^۴Leibler

آنتروپی نسبی همواره نامنفی است. همچنین $D(q||p) = 0$ اگر و تنها اگر $p = q$ باشد. از طرفی آنتروپی نسبی فاصله واقعی بین دو توزیع نیست زیرا آنتروپی نسبی مقارن نیست $(D(q||p) \neq D(p||q))$ و در نامساوی مثلثی صدق نمی‌کند [۱۲].

اطلاعات متقابل اندازه‌گیری مقدار اطلاعاتی است که یک متغیر تصادفی درباره متغیر تصادفی دیگر دارد. همچنین کاهش عدم قطعیت یک متغیر تصادفی به دلیل آگاهی از اطلاعات متغیر تصادفی دیگر است.

تعریف ۲۰۱۰۵. متغیرهای تصادفی X و Y با تابع چگالی احتمال توأم $p(x, y)$ و توابع چگالی احتمال حاشیه‌ای $p(x)$ و $p(y)$ را در نظر بگیرید. آنتروپی نسبی بین تابع چگالی احتمال توأم و حاصل ضرب $p(x)p(y)$ را اطلاعات متقابل می‌نامیم و به صورت زیر محاسبه می‌کنیم

$$\begin{aligned} I(X, Y) &= D(p(x, y) || p(x)p(y)) \\ &= \sum_x \sum_y p(x, y) \log_2 \left(\frac{p(x, y)}{p(x)p(y)} \right) \\ &= E \left[\log_2 \left(\frac{p(x, y)}{p(x)p(y)} \right) \right]. \end{aligned} \quad (205)$$

در حالت گسسته، اطلاعات متقابل مانند آنتروپی نسبی همواره نامنفی است. این معیار در واقع میزان وابستگی دو متغیر تصادفی را نشان می‌دهد. از طرفی $I(X, Y) = 0$ اگر و تنها اگر متغیرهای تصادفی X و Y مستقل باشند. همچنین اطلاعات متقابل مقارن است. یعنی

$$I(X, Y) = I(Y, X).$$

۲۰۵ آزمون نیکویی برازش

فرض کنید نتیجه یک آزمایش تصادفی به یکی از دسته‌های دو به دو مجزای $A_1, A_2, \dots, A_k$ تعلق داشته باشد و احتمال قرار گرفتن نتیجه آزمایش تصادفی در دسته A_i برابر q_i باشد که

$$0 \leq q_i \leq 1, \quad \sum_{i=1}^k q_i = 1.$$

این توزیع احتمال مجهول را با $q(x)$ نمایش می‌دهیم. می‌خواهیم بررسی کنیم که آیا این آزمایش تصادفی دارای توزیع $p(x)$ تحت شرایط

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

می‌باشد یا خیر. به عبارت دیگر می‌خواهیم برای هر x فرض

$$H_0 : q(x) = p(x),$$

را در مقابل فرض

$$H_1 : q(x) \neq p(x),$$

برای حداقل یک مقدار x آزمون کنیم. لذا آزمایش تصادفی را n بار تکرار می‌کنیم و فرض می‌کنیم O_i فراوانی حالت‌هایی باشد که نتیجه آزمایش به A_i تعلق داشته باشد. اگر فرض H_0 برقرار باشد آنگاه

$$e_i = E(O_i) = np_i.$$

برای بررسی فرضیه فوق می‌توانیم از آماره‌های زیر استفاده کنیم.

(۱) آماره پیرسن:

آماره پیرسن به صورت زیر می‌باشد

$$X^2 = \sum_{i=1}^k \frac{(O_i - e_i)^2}{e_i}.$$

برای نمونه‌های با اندازه بزرگ X^2 دارای توزیع کای‌دو با درجه آزادی $k - 1$ می‌باشد. بنابراین اگر

$$X^2 > \chi^2(k - 1, \alpha),$$

باشد آنگاه فرض H_0 رد می‌شود که در آن $\chi^2(k - 1, \alpha)$ چندک $1 - \alpha$ از توزیع کای‌دو است [۲۱].

(۲) آماره نسبت درستنمایی:

آماره نسبت درستنمایی به صورت زیر می‌باشد

$$G^2 = 2 \sum_{i=1}^k O_i \log_2 \left(\frac{O_i}{e_i} \right).$$

برای نمونه‌های با اندازه بزرگ G^2 دارای توزیع کای‌دو با درجه آزادی $k - 1$ می‌باشد. بنابراین اگر

$$G^2 > \chi^2(k - 1, \alpha),$$

باشد آنگاه فرض H_0 رد می‌شود. آماره G^2 به صورت مجانبی با آماره پیرسن X^2 معادل است [۱۲].

(۳) آماره کولموگروف-اسمیرنوف:

آماره کولموگروف-اسمیرنوف به صورت زیر می‌باشد

$$T_n = \sup_{x \in S(x)} \sqrt{n} |F_n(x) - F(x)|,$$

که در آن $F_n(x)$ توزیع تجربی از نمونه و $F(x)$ تابع توزیع متغیر تصادفی پیوسته X است. اگر T_n از مقدار بحرانی جداول مربوطه بیشتر باشد آنگاه فرض H_0 رد می‌شود [۲۱].

۳.۵ آزمون نیکویی برازش مبتنی بر آنتروپی نسبی

در این بخش با استفاده از آنتروپی نسبی یک آماره جدید برای آزمون نیکویی برازش معرفی می‌کنیم. فرض $H_0 : q = p$ را در مقابل $H_1 : q \neq p$ در نظر بگیرید که با فرضیه زیر معادل است

$$H_0 : D(q||p) = 0$$

$$H_1 : D(q||p) > 0$$

که در آن

$$D(q||p) = \sum_{i=1}^k q_i \log_2 \left(\frac{q_i}{p_i} \right) = E \left[\log_2 \left(\frac{q}{p} \right) \right].$$

فرض کنید O_i فراوانی حالت‌هایی باشد که نتیجه آزمایش به A_i تعلق دارد و $e_i = np_i$ باشد. برآوردگر درستنمایی ماکسیم q_i به صورت $\hat{q}_i = \frac{O_i}{n}$ است. چون $p_i = \frac{e_i}{n}$ می‌باشد لذا برآوردگر درستنمایی ماکسیم برای $D(q||p)$ برابر است با:

$$\hat{D}(q||p) = \frac{1}{n} \sum_{i=1}^k O_i \log_2 \left(\frac{O_i}{e_i} \right). \quad (3.5)$$

می‌دانیم که $O = (O_1, O_2, \dots, O_k)$ دارای توزیع چندجمله‌ای است. به عبارت دیگر

$$O = (O_1, O_2, \dots, O_k) \sim M_k(n, (q_1, q_2, \dots, q_k)).$$

برای یک نمونه با اندازه بزرگ O دارای توزیع نرمال چندمتغیره متقارن $N_k(nq, n(D_q - qq'))$ است که در آن D_q یک ماتریس قطری با عناصر قطری q_i می‌باشد و $q = (q_1, q_2, \dots, q_n)$ است. در نتیجه

$$\sqrt{n} \left(\frac{1}{n} O - q \right) \longrightarrow N_k(0, D_q - qq').$$

با استفاده از (۳.۵) و قضایای حد داریم

$$Z = \sqrt{n} \left(\frac{\hat{D}(q||p) - D(q||p)}{\hat{\sigma}} \right) \longrightarrow N(0, 1), \quad (4.5)$$

که در آن

$$\hat{\sigma}^2 = \frac{1}{n} \left[\sum_i O_i \left(\log_2 \frac{O_i}{e_i} \right)^2 - \left(\sum_i O_i \log_2 \frac{O_i}{e_i} \right)^2 \right]. \quad (5.5)$$

اکنون آماره زیر را در نظر می‌گیریم

$$Z_0 = \frac{\sqrt{n}\hat{D}(q||p)}{\hat{\sigma}}. \quad (۶.۵)$$

اگر $Z_0 > Z_\alpha$ باشد فرض H_0 را رد می‌کنیم که در آن Z_α چندک $1 - \alpha$ از توزیع نرمال استاندارد می‌باشد.

تذکر ۱.۳.۵. آماره Z_0 به آماره G^2 بسیار نزدیک است اما حساسیت آن از G^2 بیشتر می‌باشد. از طرفی نمی‌توانیم آماره جدید را با آزمون نیکویی برازش مبتنی بر آماره‌های X^2 ، G^2 و T_n مقایسه کنیم زیرا آماره جدید برای نمونه‌های با اندازه بسیار بزرگ مناسب است و دیگر نمی‌توان آن را به صورت یک مدل پارامتری توصیف کرد [۲۱]. لذا در ادامه از مطالعات شبیه‌سازی برای مقایسه این آماره‌ها استفاده می‌کنیم و حساسیت آن‌ها را مقایسه می‌کنیم.

۴.۵ مثال‌ها

در این بخش با استفاده از چند مثال به بررسی روش‌های معرفی شده در بخش قبل می‌پردازیم.

مثال ۱.۴.۵. از بین یک جمعیت بزرگ نمونه‌ای ۲۰۰ نفری را انتخاب کرده و تعداد دفعات مراجعه آن‌ها به یک شرکت بیمه در دوره ۴ ساله را در نظر می‌گیریم. نتایج حاصل در جدول ۱۰.۵ ارائه شده است. فرض کنید X تعداد دفعاتی باشد

جدول ۱۰.۵: دفعات مراجعه نمونه ۲۰۰ نفری به شرکت بیمه

شماره	۰	۱	۲	۳	۴	۵	۶	۷
فراوانی	۲۲	۵۳	۵۸	۳۹	۲۰	۵	۲	۱

که یک شخص به شرکت بیمه مراجعه می‌کند. می‌خواهیم فرضیه زیر را بررسی کنیم

$$H_0 : X \sim P(\lambda)$$

$$H_1 : X \not\sim P(\lambda)$$

آماره آزمون $Z_0 = ۰.۸۰۱۲$ و $Z_{۰.۰۵} = ۱.۶۴۵$ است. چون $Z_{۰.۰۵} > Z_0$ است لذا فرض H_0 رد نمی‌شود. همچنین آماره پیرسون $X^2 = ۲۳$ و $\chi^2(۴, ۰.۰۵) = ۹.۴۸۷$ می‌باشد و در این آزمون نیز H_0 رد نمی‌شود.

مثال ۲.۴.۵. نتایج یک نظرسنجی درباره موضوعی خاص در یک کشور به صورت جدول ۲۰.۵ است. یک نمونه صدتایی از این جامعه انتخاب شده است که نتایج حاصل در جدول ۳۰.۵ ارائه شده است. آیا نتایج حاصل در جدول ۳۰.۵ با نتایج نظرسنجی در جامعه مطابقت دارد؟

جدول ۲.۵: نتایج نظرسنجی در جامعه

نظر	خیلی موافق	موافق	ممتنع	مخالف	خیلی مخالف
درصد	۲۰	۳۰	۲۰	۲۰	۱۰

جدول ۳.۵: نتایج نظرسنجی در نمونه ۱۰۰ تایی

نظر	خیلی موافق	موافق	ممتنع	مخالف	خیلی مخالف
فراوانی	۱۴	۱۸	۱۸	۲۶	۲۴

با استفاده از (۶.۵) آماره $Z = ۲/۳۱۵۶$ حاصل می شود. همچنین از جدول توزیع نرمال مقدار $Z_{۰/۰۵} = ۱/۶۴۵$ به دست می آید. با مقایسه این مقادیر $Z_{۰/۰۵} > Z_۰$ است. در نتیجه فرض $H_۰$ رد می شود و نتایج حاصل در نمونه ۱۰۰ تایی با نتایج کلی جامعه مطابقت ندارد. همچنین آماره پیرسن به صورت $X^2 = ۴۶/۵$ و $\chi^2(۴, ۰/۰۵) = ۹/۴۸۷$ در سطح معنی داری $\alpha = ۰/۰۵$ است. در این حالت نیز فرض $H_۰$ رد می شود.

مثال ۳.۴.۵. یک تاس ۱۲۰ بار پرتاب و نتایج حاصل در جدول ۴.۵ ارائه شده است.

جدول ۴.۵: نتایج حاصل از ۱۲۰ بار پرتاب یک تاس

تاس	۱	۲	۳	۴	۵	۶
فراوانی	۱۸	۲۳	۱۶	۲۱	۱۸	۲۴

می دانیم که

$$p_i = \frac{1}{6} \quad e_i = 20 \times \frac{1}{6} = 20 \quad i = 1, 2, \dots, 6.$$

در نتیجه آماره $Z = ۰/۷۹۱۴$ حاصل می شود. از جدول توزیع نرمال نیز $Z_{۰/۰۵} = ۱/۶۴۵$ است. چون $Z_{۰/۰۵} > Z_۰$ می باشد لذا فرض $H_۰$ رد نمی شود. همچنین آماره پیرسن به صورت $X^2 = ۲/۵$ و $\chi^2(۵, ۰/۰۵) = ۱۱/۱$ می باشد که در این حالت نیز فرض $H_۰$ رد نمی شود.

۵.۵ شبیه سازی

کلاس فرضیه های H_1 خیلی بزرگ است و همه توزیع ها به جز توزیع های $H_۰$ را شامل می شود. لذا نمی توانیم تابع توان آماره جدید را با روش های معمول محاسبه و مقایسه کنیم. در این بخش با استفاده از توزیع های پواسن و نمایی اعداد تصادفی را تولید کرده و به کمک آن ها آزمون نیکویی برازش را انجام می دهیم. برای مقایسه توابع آزمون از مقادیر p

استفاده می‌کنیم.

۱.۵.۵ آزمون پواسن

در آزمون پواسن ۱۰۰۰ عدد تصادفی از توزیع $P(\lambda_1)$ را به‌ازای $\lambda_1 = 1, 2$ شبیه‌سازی کرده و فرضیه‌های زیر را در نظر می‌گیریم

$$H_0 : X \sim P(\lambda_2)$$

$$H_1 : X \not\sim P(\lambda_2)$$

که مقادیر λ_2 در جدول‌های ۵.۵ و ۶.۵ آمده است. همچنین

$$\hat{d} = \hat{D}(q||p), \quad \hat{\sigma}^2 = S_d^2.$$

داده‌های ارائه شده در جدول ۵.۵ از توزیع پواسن به‌ازای $\lambda = 1$ تولید شده‌اند. در سطرهای I مقادیر λ_2 نزدیک یک هستند و فرض H_0 پذیرفته می‌شود. اما در سطرهای III مقادیر λ_2 از مقدار یک فاصله دارند و فرض H_0 رد می‌شود. اگر مقادیر λ_2 بین ۱/۰۱۹۶ و ۱/۰۲ باشند آنگاه آماره جدید مبتنی بر توزیع Z فرض H_0 را به درستی رد می‌کند ولی آماره‌های X^2 و G^2 فرض H_0 را به اشتباه قبول می‌کنند. در نتیجه آماره جدید در مقایسه با آماره‌های موجود در رد فرض H_0 حساس‌تر است. به‌طور مشابه داده‌های ارائه شده در جدول ۶.۵ از توزیع پواسن به‌ازای $\lambda = 2$ تولید شده‌اند. در این حالت به‌ازای $\lambda_2 \in (2/۰۰۵, 2/۰۳)$ آماره جدید فرض H_0 را رد و آماره‌های دیگر فرض H_0 را قبول می‌کنند.

جدول ۵.۵: نتایج آزمون پواسن برای ۱۰۰۰ عدد تصادفی تولید شده با $P(1)$								
$p(G^2)$	$p(X^2)$	$p(Z)$	G^2	X^2	Z	S_d^2	$\hat{d}$	λ_2
۰/۸۸	۰/۸۹	۰/۱۲۵۱	۲/۳۶	۲/۳۲	۱/۱۵	۰/۰۰۴۸	۰/۰۰۰۸	۱
۰/۷۷۹	۰/۷۷۸۸	۰/۱۰۷۵	۲/۲۹	۲/۲۱	۱/۲۴	۰/۰۰۴۹	۰/۰۰۰۸	۱/۰۰۵
۰/۵۸۴	۰/۶	۰/۰۸۵۵	۲/۷۱	۲/۵۸	۱/۳۵	۰/۰۰۵	۰/۰۰۰۹	۱/۰۰۱
۰/۴۷	۰/۲	۰/۰۶۸	۶/۶۱	۶/۲۳	۱/۲۹	۰/۰۰۵	۰/۰۰۱۱	۱/۰۱۵
۰/۲۱۵	۰/۲۲	۰/۰۵۲۶	۸/۲۸	۸/۲۵	۱/۶۳	۰/۰۰۶	۰/۰۰۱۲	۱/۰۱۹
۰/۲۰۱	۰/۲۱	۰/۰۵۰۵	۸/۷۲	۸/۴۹	۱/۶۲۶	۰/۰۰۶	۰/۰۰۱۳	۱/۰۱۹۵
۰/۱۹۸	۰/۲۱	۰/۰۴۹۵	۸/۷۹	۸/۵۴	۱/۶۵	۰/۰۰۶	۰/۰۰۱۳	۱/۰۱۹۶
۰/۱۹۶	۰/۲۰۹	۰/۰۴۹۵	۸/۸۴	۸/۵۹	۱/۶۵۴	۰/۰۰۶	۰/۰۰۱۳	۱/۰۱۹۷
۰/۱۸۸	۰/۲	۰/۰۴۸۵	۸/۹۹	۸/۷۵	۱/۶۶	۰/۰۰۶	۰/۰۰۱۳	۱/۰۲
۰/۰۲	۰/۰۲۳	۰/۰۲۰۲	۱۵/۲	۱۲/۷۷	۲/۰۵	۰/۰۰۷	۰/۰۰۱۷	۱/۰۳
۰	۰	۰/۰۰۶۲	۲۲/۲۷	۲۲/۶	۲/۵	۰/۰۰۸۵	۰/۰۰۴۳	۱/۰۴
۰	۰	۰/۰۰۱۵	۳۲/۱۸	۳۲/۱۸	۲/۹۶	۰/۰۰۱	۰/۰۰۰۳	۱/۰۵

جدول ۶.۵: نتایج آزمون پواسن برای ۱۰۰۰ عدد تصادفی تولید شده با $P(2)$

$p(G^2)$	$p(X^2)$	$p(Z)$	G^2	X^2	Z	S_d^2	$\hat{d}$	λ_2	
۰/۸۶	۰/۸۸	۰/۰۵۷۱	۳/۸۱	۳/۷۲	۱/۵۸	۰/۰۰۵	۰/۰۰۱	۲	I
۰/۸۵۵	۰/۸۶	۰/۰۵۴۸	۳/۹۶	۳/۸۷	۱/۶	۰/۰۰۵	۰/۰۰۱	۲/۰۰۱	
۰/۷۹	۰/۸	۰/۰۴۹۵	۴/۶۴	۴/۵۳	۱/۶۵	۰/۰۰۵	۰/۰۰۱۷	۲/۰۰۵	
۰/۶۸	۰/۶۹	۰/۰۴۱۸	۵/۷۲	۵/۵۸	۱/۷۳	۰/۰۰۵	۰/۰۰۱۲	۲/۰۰۱	II
۰/۱۴	۰/۱۶	۰/۰۱۵۰	۱۲/۴۷	۱۲/۱۷	۲/۱۷	۰/۰۰۶	۰/۰۰۲	۲/۰۰۳	
۰/۰۲۸	۰/۰۳۳	۰/۰۰۷۱	۱۷/۳	۱۶/۸۹	۲/۴۵	۰/۰۰۸	۰/۰۰۲	۲/۰۰۴	
۰	۰	۰/۰۰۳	۲۳/۰۸	۲۲/۵۲	۲/۷۵	۰/۰۰۸	۰/۰۰۳	۲/۰۰۵	III
۰	۰	۰/۰۰۱۱	۲۹/۸۲	۲۹/۰۷	۳/۰۶	۰/۰۰۹۵	۰/۰۰۳	۲/۰۰۶	
۰	۰	۰	۳۷/۴۸	۳۶/۵۲	۳/۳۸	۰/۰۱۱	۰/۰۰۳۵	۲/۰۰۷	

۲.۵.۵ آزمون نمایی

مشابه حالت قبل تعداد ۱۰۰۰ عدد تصادفی از توزیع نمایی $\exp(\lambda_1)$ را به ازای $\lambda_1 = 1$ شبیه سازی کرده و فرضیه های زیر را در نظر می گیریم

$$H_0 : X \sim \exp(\lambda_2)$$

$$H_1 : X \not\sim \exp(\lambda_2)$$

که مقادیر λ_2 در جدول ۷.۵ آمده است. همچنین مقادیر $\hat{d}$ ، $\hat{\sigma}^2$ و مقادیر p وابسته به آماره های Z ، X^2 ، G^2 و T_n محاسبه شده است. از نتایج ارائه شده در جدول ۷.۵ واضح است که به ازای مقادیر λ_2 نزدیک یک در سطر IV فرض H_0 پذیرفته می شود و به ازای $\lambda_2 = 1/06$ در سطرهای I که دور از یک می باشد فرض H_0 رد می شود. همچنین به ازای $\lambda \in (1/039, 1/041)$ آماره جدید Z فرض H_0 را رد می کند در حالی که آماره های X^2 ، G^2 و T_n فرض H_0 را می پذیرند. در نتیجه آماره جدید در مقایسه با آماره های دیگر حساس تر است.

۶.۵ آزمون مستقل بودن مبتنی بر اطلاعات متقابل

یک نمونه تصادفی به اندازه n از یک جامعه را در نظر بگیرید. مشاهدات در نمونه تصادفی براساس دو معیار طبقه بندی می شود. با استفاده از معیار اول هر مشاهده با یکی از ردیف ها (R) و با استفاده از معیار دوم هر مشاهده با یکی از ستون ها (C) مرتبط می شود. فرض کنید O_{ij} تعداد مشاهدات مرتبط با سطر i ام و ستون j ام باشد. این مقادیر را در جدول احتمالی از مرتبه $R \times C$ به صورت زیر قرار می دهیم.

جدول ۷.۵: نتایج آزمون نمایی برای ۱۰۰۰ عدد تصادفی تولید شده با $\exp(1)$

$p(G^\dagger)$	$p(X^\dagger)$	$p(Z)$	T_n	$G^\dagger$	$X^\dagger$	Z	$S_d^\dagger$	$\hat{d}$	λ_2	
۰/۷۶۶	۰/۷۶۷	۰/۱۲۵۱	۰/۷۶	۴/۱	۴/۰۹	۱/۱۵	۰/۰۰۱۲	۰/۰۰۰۴	۱	I
۰/۷	۰/۶۷۳	۰/۱۲۱	۰/۴۰	۴/۶۶	۴/۹	۱/۱۷	۰/۰۰۱۴	۰/۰۰۰۴	۱/۰۲	
۰/۳۸	۰/۳۳۷	۰/۰۷۷۸	۰/۶۴	۷/۶۲	۸/۱	۱/۴۲	۰/۰۰۲۱	۰/۰۰۰۷	۱/۰۳	
۰/۱۱	۰/۰۸۶۲	۰/۰۴۲۷	۰/۸۷	۱۱/۷۵	۱۲/۵۸	۱/۷۲	۰/۰۰۰۳	۰/۰۰۰۹	۱/۰۳۹	II
۰/۰۹	۰/۰۷۲	۰/۰۴۰۱	۰/۹۰	۱۲/۳	۱۳/۱۸	۱/۷۵	۰/۰۰۳۲	۰/۰۰۰۱	۱/۰۴	
۰/۰۸۳	۰/۰۵۷۳	۰/۰۳۶۷	۰/۹۲	۱۲/۸۶	۱۳/۷۹	۱/۷۹	۰/۰۰۳۳	۰/۰۰۰۱	۱/۰۴۱	
۰/۰۰۹	۰/۰۰۵۵	۰/۰۱۶۶	۱/۱۰	۱۸/۶۶	۲۰/۱۳	۲/۱۳	۰/۰۰۴۶	۰/۰۰۱۴	۱/۰۵	III
۰	۰	۰/۰۰۵۹	۱/۳۸	۲۶/۶۷	۲۸/۹۷	۲/۵۲	۰/۰۰۶۵	۰/۰۰۰۲	۱/۰۶	IV

Y	۱	۲	...	C	Sum
X					
۱	O_{11}	O_{12}	...	O_{1C}	$O_{1\cdot}$
۲	O_{21}	O_{22}	...	O_{2C}	$O_{2\cdot}$
$\vdots$	$\vdots$	$\vdots$	$\ddots$	$\vdots$	$\vdots$
R	O_{R1}	O_{R2}	...	O_{RC}	$O_{R\cdot}$
Sum	$O_{\cdot 1}$	$O_{\cdot 2}$	...	$O_{\cdot C}$	n

مجموع مقادیر O_{ij} در سطر i برابر $O_{i\cdot}$ ، مجموع مقادیر O_{ij} در سطر j برابر $O_{\cdot j}$ و مجموع کل آن‌ها برابر n است. فرض H_0 را مستقل بودن یک مشاهده در سطر i ام از برخی مشاهدات در ستون j ام در نظر می‌گیریم. بنابراین فرض H_0 :

$$p(i \text{ سطر}, j \text{ ستون}) = p(i \text{ سطر}) \times p(j \text{ ستون}).$$

فرض H_1 :

$$p(i \text{ سطر}, j \text{ ستون}) \neq p(i \text{ سطر}) \times p(j \text{ ستون}).$$

آماره آزمون به صورت

$$X^2 = \sum_j \sum_i \frac{(O_{ij} - e_{ij})^2}{e_{ij}},$$

می‌باشد که در آن

$$e_{ij} = \frac{O_{i\cdot} \times O_{\cdot j}}{n}.$$

این آماره دارای توزیع کای دو با درجه آزادی $(R-1)(C-1)$ است. لذا اگر $\chi^2 > \chi^2((R-1)(C-1), \alpha)$ باشد آن‌گاه فرض H_0 رد می‌شود. با استفاده از تعریف اطلاعات متقابل فرض‌های H_0 و H_1 را می‌توانیم به صورت زیر بیان کنیم

$$H_0 : I(X, Y) = 0$$

$$H_1 : I(X, Y) \neq 0$$

یا به طور معادل

$$H_0 : D(p(x, y) || p(x)p(y)) = 0$$

$$H_1 : D(p(x, y) || p(x)p(y)) > 0$$

همچنین برآوردگر ماکسیم درستمایی برای $I(X, Y)$ برابر است با:

$$\hat{I}(X, Y) = \frac{1}{n} \sum_i \sum_j O_{ij} \log_2 \left(\frac{O_{ij}}{e_{ij}} \right),$$

که در آن $O = (O_{11}, O_{12}, \dots, O_{RC})$ مقادیر مشاهده شده است که دارای توزیع چندجمله‌ای می‌باشد و اگر اندازه نمونه به اندازه کافی بزرگ باشد دارای توزیع نورمال خواهد بود. به عبارت دیگر

$$Z = \frac{\sqrt{n}(\hat{I}(X, Y) - I(X, Y))}{\hat{\sigma}} \rightarrow N(0, 1),$$

که در آن

$$\hat{\sigma}^2 = \frac{1}{n} \left[\sum_i \sum_j O_{ij} \left(\log_2 \frac{O_{ij}}{e_{ij}} \right)^2 - \left(\sum_i \sum_j O_{ij} \log_2 \frac{O_{ij}}{e_{ij}} \right)^2 \right].$$

سرانجام اگر

$$Z_0 = \frac{\sqrt{n}(\hat{I}(X, Y))}{\hat{\sigma}} > Z_\alpha,$$

آن‌گاه فرض H_0 رد می‌شود.

مثال ۱.۶.۵. تعداد ۱۲۰۰ نفر در چهار گروه و دو دیدگاه متفاوت طبقه‌بندی شده‌اند که نتایج حاصل در جدول ۸.۵ ارائه شده است. آیا این دو نوع دیدگاه مستقل از هم هستند؟

با استفاده از داده‌های جدول ۸.۵ آماره $\chi^2 = ۲۰/۵۹$ و از جدول توزیع خی دو $\chi^2(3, 0/05) = ۷/۸۱۵$ حاصل می‌شود. چون $\chi^2 > \chi^2(3, 0/05)$ است لذا فرض مستقل بودن رد می‌شود. همچنین $Z_0 = ۲/۵۱۴$ و $Z_0 = ۱/۶۴۵$ است و در نتیجه فرض مستقل بودن نیز با این آماره رد می‌شود. برای این مثال هر دو آماره نتیجه یکسانی ارائه می‌دهند.

جدول ۸.۵: ۱۲۰۰ نفر در چهار گروه با دو دیدگاه

دیدگاه گروه	I	II	مجموع
A	۳۲	۲۶۹	۳۰۰
B	۵۱	۱۹۹	۲۵۰
C	۶۷	۲۳۳	۳۰۰
D	۸۳	۲۶۷	۳۵۰
مجموع	۲۳۳	۹۶۷	۱۲۰۰

مثال ۲.۶.۵. برای بررسی تأثیر یک داروی شیمیایی روی بذر، ۲۵۰ بذر را انتخاب می‌کنیم که برای ۱۰۰ مورد از آن‌ها این دارو استفاده شده است و برای ۱۵۰ بذر دیگر از این دارو استفاده نشده است. نتایج حاصل در جدول ۹.۵ ارائه شده است.

جدول ۹.۵: تأثیر داروی شیمیایی روی بذر

	مثبت	منفی
با داروی شیمیایی	۸۴	۱۶
بدون داروی شیمیایی	۱۳۲	۱۸

با استفاده از داده‌های جدول ۹.۵ آماره $X^2 = ۰/۸۱۷$ و از جدول توزیع خی‌دو $\chi^2(۱, ۰/۰۵) = ۳/۸۴۱$ حاصل شده است. چون $\chi^2(۳, ۰/۰۵) > X^2$ است لذا فرض مستقل بودن قبول می‌شود. همچنین آماره جدید $Z_{۰/۰۵} = ۰/۴۴۵۶$ و $Z_{۰/۰۵} = ۱/۶۴۵$ است و در نتیجه فرض مستقل بودن نیز با این آماره رد نمی‌شود.