

Package ‘chisquare’

October 29, 2025

Title Chi-Square and G-Square Test of Independence, Power and Residual Analysis, Measures of Categorical Association

Version 1.2

Description Provides the facility to perform the chi-square and G-square test of independence, calculates the retrospective power of the traditional chi-square test, compute permutation and Monte Carlo p-value, and provides measures of association for tables of any size such as Phi, Phi corrected, odds ratio with 95 percent CI and p-value, Yule' Q and Y, adjusted contingency coefficient, Cramer's V, V corrected, V standardised, bias-corrected V, W, Cohen's w, Goodman-Kruskal's lambda, and tau. It also calculates standardised, moment-corrected standardised, and adjusted standardised residuals, and their significance, as well as the Quetelet Index, IJ association factor, and adjusted standardised counts. It also computes the chi-square-maximising version of the input table. Different outputs are returned in nicely formatted tables.

Depends R (>= 4.0)

Imports graphics, gt (>= 0.3.1), stats

License GPL (>= 2)

Encoding UTF-8

LazyData true

RoxygenNote 7.2.3

NeedsCompilation no

Author Gianmarco Alberti [aut, cre]

Maintainer Gianmarco Alberti <gianmarcoalberti@gmail.com>

Repository CRAN

Date/Publication 2025-10-29 09:00:08 UTC

Contents

chisquare	2
diseases	26
safety	27
social_class	27
Index	28

chisquare	<i>R function for Chi-square, (N-1) Chi-square, and G-Square test of independence, power calculation, measures of association, and standardised/moment-corrected standardised/adjusted standardised residuals, visualisation of odds ratio in 2xk tables (where k >= 2)</i>
-----------	--

Description

The function performs the chi-square test (both in its original format and in the N-1 version) and the G-square test of independence on the input contingency table. It also calculates the retrospective power of the traditional chi-square test and various measures of categorical association for tables of any size, returns standardised, moment-corrected standardised, adjusted standardised residuals (with indication of their significance), Quetelet Index, IJ association factor, adjusted standardised counts, and the chi-square-maximising version of the input table. It also calculates relative and absolute contributions to the chi-square statistic. The p-value associated with the chi-square statistic is also computed via both permutation-based and Monte Carlo methods, using the exact p-value calculation method of Phipson & Smyth (2010), which ensures that p-values are never zero and that Type I error rates are correctly controlled. Nicely-formatted output tables are rendered. Optionally, in 2xk tables (where k >= 2), a plot of the odds ratios can be rendered.

Visit this [LINK](#) to access the package's vignette.

Usage

```
chisquare(
  data,
  B = 1000,
  plot.or = FALSE,
  reference.level = 1,
  row.level = 1,
  or.alpha = 0.05,
  power.alpha = 0.05,
  adj.alpha = FALSE,
  marginal.type = "average",
  custom.row.totals = NULL,
  custom.col.totals = NULL,
  format = "short",
  render.all.tables = FALSE,
  graph = FALSE,
  oneplot = TRUE,
  tfs = 13
)
```

Arguments

data	Dataframe containing the input contingency table.
------	---

B	Number of resamples for all simulation-based computations. This parameter controls: (1) the number of permutations for the permutation-based chi-square test, (2) the number of Monte Carlo simulations for the Monte Carlo chi-square test, and (3) the number of bootstrap samples for computing confidence intervals around the W coefficient and PEM values. Default is 999, which gives minimum p-value = 0.001, p-value resolution = 0.001, and stable bootstrap percentiles. For publication-quality results, users might want to use B = 9999 (minimum p-value = 0.0001). P-values are calculated using the exact method of Phipson & Smyth (2010): $p = (B + 1)/(m + 1)$, where B is the count of simulated statistics at least as extreme as observed, and m is the number of simulations.
plot.or	Takes TRUE or FALSE (default) if the user wants a plot of the odds ratios to be rendered (only for 2xk tables, where $k \geq 2$).
reference.level	The index of the column reference level for odds ratio calculations (default: 1). The user must select the column level to serve as the reference level (only for 2xk tables, where $k \geq 2$).
row.level	The index of the row category to be used in odds ratio calculations (1 or 2; default: 1). The user must select the row level to which the calculation of the odds ratios make reference (only for 2xk tables, where $k \geq 2$).
or.alpha	The significance level used for the odds ratios' confidence intervals (default: 0.05).
power.alpha	The significance level used for the calculation of the power of the traditional chi-square test (default: 0.05).
adj.alpha	Takes TRUE or FALSE (default) if the user wants or does not want the significance level of the residuals (standardised, adjusted standardised, and moment-corrected) to be corrected using the Sidak's adjustment method (see Details).
marginal.type	Defines the target marginal sums used for table standardisation via Iterative Proportional Fitting. It takes <i>average</i> (default) to have target row and column marginals equal to the table's grand total divided by the number of rows and columns, respectively; it takes <i>percent</i> to have target marginals equal to fractions of a grand total set to 100. See Details.
custom.row.totals	A vector of numbers indicating the target row marginals to be used for table standardisation via Iterative Proportional Fitting (NULL by default; see Details).
custom.col.totals	A vector of numbers indicating the target column marginals to be used for table standardisation via Iterative Proportional Fitting (NULL by default; see Details).
format	Takes <i>short</i> (default) if the dataset is a dataframe storing a contingency table; if the input dataset is a dataframe storing two columns that list the levels of the two categorical variables, <i>long</i> will preliminarily cross-tabulate the levels of the categorical variable in the 1st column against the levels of the variable stored in the 2nd column.
render.all.tables	Takes TRUE or FALSE (default) if the user wants or does not want all the 'gt' tables to be automatically rendered.

graph	Takes TRUE or FALSE (default) if the user wants or does not want to plot the permutation and Monte Carlo distribution of the chi-square statistic across the number of simulated tables set by the B parameter.
oneplot	Takes TRUE (default) or FALSE if the user wants or does not want to render of the permutation and Monte Carlo distribution in the same plot.
tfs	Numerical value to set the size of the font used in the main body of the various output tables (13 by default).

Details

HYPOTHESIS TESTING AND POWER ANALYSIS

Suggestion of a suitable chi-square testing method

The first rendered table includes a suggestion for the applicable chi-squared test method, derived from an internal analysis of the input contingency table. The decision logic used is as follows:

For 2x2 tables:

- if the grand total is equal to or larger than 5 times the number of cells, the traditional Chi-Square test is suggested. Permutation or Monte Carlo methods can also be considered.

- if the grand total is smaller than 5 times the number of cells, the minimum expected count is checked:

- (A) if it is equal to or larger than 1, the $(N-1)/N$ adjusted Chi-Square test is suggested, with an option for Permutation or Monte Carlo methods.

- (B) if it is less than 1, the Permutation or Monte Carlo method is recommended.

For tables larger than 2x2:

- the logic is similar to that for 2x2 tables, with the same criteria for suggesting the traditional Chi-Square test, the $(N-1)/N$ adjusted test, or the Permutation or Monte Carlo methods.

The rationale of a threshold for the applicability of the traditional chi-square test corresponding to 5 times the number of cells is based on the following.

Literature indicates that the traditional chi-squared test's validity is not as fragile as once thought, especially when considering the average expected frequency across all cells in the cross-tab, rather than the minimum expected value in any single cell. An average expected frequency of at least 5 across all cells of the input table should be sufficient for maintaining the chi-square test's reliability at the 0.05 significance level.

As a consequence, a table's grand total equal to or larger than 5 times the number of cells should ensure the applicability of the traditional chi-square test (at alpha 0.05).

See: Roscoe-Byars 1971; Greenwood-Nikulin 1996; Zar 2014; Alberti 2024.

For the rationale of the use of the $(N-1)/N$ adjusted version of the chi-square test, and for the permutation and Monte Carlo method, see below.

Chi-square statistics adjusted using the $(N-1)/N$ adjustment

The adjustment is done by multiplying the chi-square statistics by $(N-1)/N$, where N is the table

grand total (sample size). The p-value of the corrected statistic is calculated the regular way (i.e., using the same degrees of freedom as in the traditional test). The correction seems particularly relevant for tables where N is smaller than 20 and where the expected frequencies are equal or larger than 1. The corrected chi-square test proves more conservative when the sample size is small. As N increases, the term $(N-1)/N$ approaches 1, making the adjusted chi-square value virtually equivalent to the unadjusted value.

See: Upton 1982; Rhoades-Overall 1982; Campbel 2007; Richardson 2011; Alberti 2024.

Permutation-based and Monte Carlo p-value for the chi-square statistic

The p-value of the observed chi-square statistic is also calculated using both a permutation-based and a Monte Carlo approach. In the first case, the dataset is permuted B times (999 by default), whereas in the second method B establishes the number of random tables generated under the null hypothesis of independence (999 by default).

Important: P-values are computed using the method of Phipson & Smyth (2010), which calculates exact p-values when permutations are randomly drawn. The formula used is $p = (B + 1)/(m + 1)$, where B is the number of simulated test statistics at least as extreme as the observed statistic, and m is the total number of simulations. This approach treats the observed data as one valid arrangement under the null hypothesis, ensuring that:

- P-values are never exactly zero (minimum p-value = $1/(m + 1)$)
- Type I error rates are correctly controlled
- The test has the proper size, even when the number of simulations is relatively small

The same value of B is also used to generate bootstrap samples for computing confidence intervals around certain effect size measures (W coefficient and PEM values). This ensures computational efficiency and consistency across all simulation-based estimates.

As for the permutation method, the function does the following internally:

- (1) Converts the input dataset to long format and expands to individual observations;
- (2) Calculates the observed chi-squared statistic;
- (3) Randomly shuffles (B times) the labels of the levels of one variable, and recalculates chi-squared statistic for each shuffled dataset;
- (4) Computes the p-value based on the distribution of permuted statistics (see below).

For the rationale of the permutation-based approach, see for instance Agresti et al 2022.

For the rationale of the Monte Carlo approach, see for instance the description in Beh-Lombardo 2014: 62-64.

Both distributions can be optionally plotted setting the `graph` parameter to `TRUE`.

Retrospective Power of the Traditional Chi-Square Test

The function calculates the (highly controversial) retrospective power of the traditional chi-square test, which is the probability of correctly rejecting the null hypothesis when it is false. The power is determined by the observed chi-square statistic, the sample size, and the degrees of freedom,

without explicitly calculating an effect size, following the method described by Oyeyemi et al. 2010.

The degrees of freedom are calculated as $(\text{number of rows} - 1) * (\text{number of columns} - 1)$. The alpha level is set by default at 0.05 and can be customized using the `power.alpha` parameter. The power is then estimated using the non-centrality parameter based on the observed chi-square statistic.

The calculation involves determining the critical chi-squared value based on the alpha level and degrees of freedom, and then computing the probability that the chi-squared distribution with the given degrees of freedom exceeds this critical value.

The resulting power value indicates how likely the test is to detect an effect if one exists. A power value close to 1 suggests a high probability of detecting a true effect, while a lower value indicates a higher risk of a Type II error. Typically, a power value of 0.8 or higher is considered robust in most research contexts.

DECOMPOSITION OF THE CHI-SQUARE STATISTIC

Cells' relative contribution (in percent) to the chi-square statistic

The cells' relative contribution (in percent) to the chi-square statistic is calculated as:

$chisq.values / chisq.stat * 100$, where

$chisq.values$ and $chisq.stat$ are the chi-square value in each individual cell of the table and the value of the chi-square statistic, respectively. The *average contribution* is calculated as $100 / (nr * nc)$, where nr and nc are the number of rows and columns in the table respectively.

Cells' absolute contribution (in percent) to the chi-square statistic

The cells' absolute contribution (in percent) to the chi-square statistic is calculated as:

$chisq.values / n * 100$, where

$chisq.values$ and n are the chi-square value in each individual cell of the table and the table's grand total, respectively. The *average contribution* is calculated as sum of all the absolute contributions divided by the number of cells in the table.

For both the relative and absolute contributions to the chi-square, see: Beasley-Schumacker 1995: 90.

TABLE STANDARDISATION

Table standardisation via Iterative Proportional Fitting

The function internally standardises the input contingency table via the *Iterative Proportional Fitting* (IPF) routine, described for instance in Reynolds 1977: 32-33. The standardised table is returned and rendered by the function.

The standardisation is performed so that the rows feature the same sums, and the columns feature the same sum, while keeping the table's grand total unchanged. This removes the effect of skewed marginal distributions, while preserving the association structure of the original table.

The rationale for table standardisation is that many association measures, among which Cramer's V, are affected by the skewness in the marginal distributions. Coefficients calculated on standardised tables are comparable across tables because the impact of different marginal distributions is controlled for. Table standardisation thus serves as a preliminary step towards coefficients comparison across tables.

The target row and column marginals used in the standardisation process are set using the `marginal.type`, `custom.row.totals`, and `custom.col.totals` parameters.

Specifying target marginals for table standardisation

The parameters `marginal.type`, `custom.row.totals`, and `custom.col.totals` control the target marginal sums for the Iterative Proportional Fitting algorithm. Users should follow one of two approaches:

Approach 1: Use a predefined marginal type (recommended for standard analyses)

Set both `custom.row.totals = NULL` and `custom.col.totals = NULL`, and specify `marginal.type`:

- `marginal.type = "average"` (default): Each row marginal is set to the mean of the observed row sums, and each column marginal is set to the mean of the observed column sums. For example, in a 3×4 table with observed row sums [30, 50, 40] and column sums [20, 25, 35, 40], the target row marginals would be [40, 40, 40] and the target column marginals would be [30, 30, 30, 30].
- `marginal.type = "percent"`: Row and column marginals are set so that the grand total equals 100, with equal marginals. In a 3×4 table, each row target would be 33.33 and each column target would be 25. This is particularly useful for interpretability when comparing across tables.

Approach 2: Specify both custom marginals (for specific analytical requirements)

Provide both `custom.row.totals` and `custom.col.totals` as numeric vectors:

- `custom.row.totals`: Must be a numeric vector with length equal to the number of rows in the table. Each element specifies the target marginal sum for the corresponding row.
- `custom.col.totals`: Must be a numeric vector with length equal to the number of columns in the table. Each element specifies the target marginal sum for the corresponding column.
- When both custom marginals are provided, `marginal.type` is ignored.
- Ensure that the sum of `custom.row.totals` equals the sum of `custom.col.totals` to maintain a consistent target grand total.

Important: The function technically allows specifying custom marginals for only one dimension (with `marginal.type` determining the other), but this is not recommended as it may result in inconsistent target grand totals and ambiguous standardisation outcomes.

It goes without saying that any target marginals can be chosen in the standardisation process. The choice of target values affects the scale of the standardised counts but not their relative relationships. By setting all the row and columns sums to unity, the function follows Rosenthal-Rosnow 2008.

In the iterative procedure, the convergence to the established marginals is reached when the counts obtained in a given iteration do not differ from the ones obtained in the previous iteration by more than a threshold internally set to 0.001. The maximum number of iterations is internally set to 10,000 to prevent infinite loop; after that, the convergence is deemed as failed.

The standardised table is used as the basis for computing Cramer's V standardised (see below under *Measures of Categorical Association*) and for the adjusted standardised counts (see below under *Post-Hoc Analysis*).

For table standardisation as a preliminary step towards coefficients comparison across tables, see for example Smith 1976; Reynolds 1977; Liu 1980.

See also: Fienberg 1971; Rosenthal-Rosnow 2008.

CHI-SQUARE-MAXIMISING TABLE

Chi-square-maximising table

The chi-square-maximising table is the version of the input cross-tab that, while preserving the marginal configuration, produces the largest divergence between the observed and the expected counts and, therefore, the maximum chi-squared value achievable given the observed marginals.

The table is worked out using the routine described by Berry and colleagues. This allocation routine effectively maximises the chi-square statistic by concentrating (to the maximum extent possible given the marginals) the levels of one variable into specific levels of the other.

As Berry and colleagues have noted, there can be alternative positions for the zeros in the chi-square-maximising table, but the obtained non-zero counts are the only ones that allow the maximisation of the chi-squared statistic.

The chi-square-maximising table is used to compute Cramer's V max and C max, which are in turn used to compute Cramer's V corrected (equal to phi-corrected for 2x2 tables) and C corrected. For the latter two, see further below.

On the chi-square-maximising table, see Berry et al. 2018.

POST-HOC ANALYSIS

Standardised residuals

The standardised residuals are calculated as follows:

$$(observed - expected) / \sqrt{expected}, \text{ where}$$

observed and *expected* are the observed and expected frequencies for each cell. The standardised residuals follow approximately a standard normal distribution under the null hypothesis of independence. However, they are not truly standardised because they have a non-unit variance, which is why the adjusted standardised residuals (see below) should be preferred for formal inference.

Adjusted standardised residuals

The adjusted standardised residuals are calculated as follows:

$stand.res[i, j] / \sqrt{(1 - sr[i]/n) * (1 - sc[j]/n)}$, where

$stand.res$ is the standardised residual for cell ij , sr is the row sum for row i , sc is the column sum for column j , and n is the table grand total. The *adjusted standardised residuals* should be used in place of the standardised residuals since the latter are not truly standardised because they have a non-unit variance. The standardised residuals therefore underestimate the divergence between the observed and the expected counts. The adjusted standardised residuals (and the moment-corrected ones) correct that deficiency.

See: Haberman 1973.

Moment-corrected standardised residuals

The moment-corrected standardised residuals are calculated as follows:

$stand.res / (\sqrt{(nr - 1) * (nc - 1) / (nr * nc)})$, where

$stand.res$ is each cell's standardised residual, nr and nc are the number of rows and columns respectively.

See Garcia-Perez-Nunez-Anton 2003: 827.

Significance of the residuals

The significance of the residuals (standardised, moment-corrected standardised, and adjusted standardised) is assessed using alpha 0.05 or, optionally (by setting the parameter `adj.alpha` to TRUE), using an adjusted alpha calculated using the Sidak's method:

$alpha.adj = 1 - (1 - 0.05)^{1/(nr * nc)}$, where

nr and nc are the number of rows and columns in the table respectively. The adjusted alpha is then converted into a critical two-tailed z value.

See: Beasley-Schumacker 1995: 86, 89.

Standardised median polish residuals

These residuals are calculated by applying median polish to the log-transformed counts, which generates a resistant fit that is less affected by outliers than traditional methods. The (Pearson) residuals are standardised by dividing the difference between observed and robustly fitted values by the square root of the fitted values. This approach addresses the problems of "masking" and "swamping" that can occur with traditional residuals (Simonoff 2003). Masking occurs when two or more outliers effectively "hide" each other so they are not identified, while swamping occurs when outliers draw the model fit towards themselves enough that non-outlying cells are mistakenly identified as outliers. The median polish method is resistant to both these effects.

Importantly, unlike traditional standardised residuals, these median polish residuals should not be interpreted with reference to the normal distribution. Instead, as recommended by Simonoff (2003), values with absolute magnitude greater than 3 are generally considered extreme. In the rendered output table, cells with residuals greater than 3 are highlighted in red, while those with residuals less than -3 are highlighted in blue.

See: Mosteller-Parunak 1985; Simonoff 2003.

Adjusted standardised median polish residuals

These residuals are computed in the same manner as the standardised ones, but they are adjusted by dividing the standardised median polish residuals by adjustment factors calculated as $\sqrt{(1 - \text{row prop}) * (1 - \text{col prop})}$ for each cell (Haberman 1973), where *row prop* and *col prop* are the row and column proportions of the observed counts. This adjustment improves the comparison of residuals across different cells by taking into account their varying variances due to different positions in the table, while maintaining the resistance to outliers provided by the median polish approach.

As with the (un-adjusted) standardised median polish residuals, these adjusted residuals are not interpreted with reference to the normal distribution. The same threshold of absolute values greater than 3 is used to identify extreme values, which are highlighted in red (>3) or blue (<-3) in the output table.

See: Mosteller-Parunak 1985.

Quetelet Index and IJ association factor

The Quetelet Index expresses the relative change in probability in one variable when considering the association with the other. The sign indicates the direction of the change, that is, whether the probability increases or decreases, compared to the probability expected under the hypothesis of independence. The IJ association factor is centred around 1.0 and represents the factor by which the probability in one variable changes when considering the association with the other. A decrease in probability is indicated by a factor smaller than 1.

The Quetelet index is computed as:

$$(\text{observedfreq} / \text{expectedfreq}) - 1$$

The IJ association factor is computed as:

$$\text{observedfreq} / \text{expectedfreq}$$

The thresholds for an IJ factor indicating a noteworthy change in probability, based on Agresti 2013, are: "larger than 2.0" and "smaller than 0.5". The thresholds for the Quetelet Index are based on those, and translate into "larger than 1.0" and "smaller than -0.50". For example, a Quetelet index greater than 1.0 indicates that the observed probability is more than double the expected probability under independence, corresponding to an increase of over 100 percent. Similarly, a Quetelet index less than -0.5 indicates that the observed probability is less than half the expected probability, corresponding to a substantial decrease of over 50 percent.

For the Quetelet Index, see: Mirkin 2023.

For the IJ association factor, see: Agresti 2013; Fagerland et al 2017. Note that the IJ association factor is called 'Pearson ratio' by Goodman 1996.

Percentage of Maximum Deviation (PEM)

The PEM is a standardised measure of local association in contingency tables that is calculated separately for each cell. For any given cell, it quantifies how far the observed frequency deviates from the expected frequency under independence, relative to the maximum possible deviation for that cell given the marginal constraints.

For each cell, when the observed frequency exceeds the expected (positive association), PEM equals $(\text{observed} - \text{expected}) / (\text{minimum of row or column total} - \text{expected}) * 100$. When observed is less than expected (negative association), PEM equals $(\text{observed} - \text{expected}) / \text{expected} * 100$. Values range from -100 indicates independence, positive values indicate attraction between categories, and negative values indicate repulsion.

Following established guidelines, PEM values under 5 values over 10 exceptionally strong relationships. The PEM's key advantages include its standardised scale enabling comparisons across cells, its consideration of marginal constraints, and its intuitive percentage interpretation.

Confidence intervals for each cell's PEM are computed using basic bootstrap resampling. The procedure generates B resampled contingency tables (with B set by default to 999, but can be customised by the user) by drawing with replacement from the original data, calculates the PEM for each cell in each resampled table, and uses the empirical percentiles of the bootstrap distribution to construct confidence intervals. Specifically, for a 95perc confidence interval, the 2.5th and 97.5th percentiles of the bootstrap distribution are used as the lower and upper bounds. This approach provides a measure of the sampling variability in the PEM estimates without making distributional assumptions.

See: Cibois 1993; Lefèvre-Champely 2009.

Adjusted standardised counts

The function computes adjusted standardised counts for the input contingency table. It first standardises the counts via Iterative Proportional Fitting (see above under *Table Standardisation*) so that all row and column totals equal unity; then adjusts these standardised counts by subtracting the table's grand mean (that is, the grand total divided by the number of cells).

The resulting adjusted standardised counts isolate the core association structure of the contingency table by removing the confounding effects of marginal distributions. This allows for meaningful comparisons across different tables, even when the original tables have different marginal totals.

In the adjusted standardised counts, a value of 0 indicates independence between the row and column variables for that cell. Positive values indicate a higher association than expected under independence, while negative values indicate a lower association than expected.

Unlike standardised residuals, these adjusted standardised counts can be directly compared across tables of the same size, making them particularly useful for comparative analyses across different samples, time periods, or contexts.

Goodman-Kruskal residuals

These residuals help identify how categories of one variable influence the predictability of categories of the other variable. Essentially, they help understand the association structure in the input cross-tab in terms of predictability. Two tables are rendered (and also returned): one considering the column variable as independent, one considering the row variable as independent. These residuals nicely complement the use of Goodman-Kruskal lambda as global measure of association with a Proportional Reduction of Error flavour.

For more details on these residuals, see: Kroonenberg-Lombardo 2010; Beh-Lombardo 2021.

MEASURES OF CATEGORICAL ASSOCIATION

The function produces the following measures of categorical association, organised by their underlying statistical foundation:

Chi-square-based measures:

- Phi (with indication of the magnitude of the effect size; only for 2x2 tables)
- Phi max (used to compute Phi corrected; only for 2x2 tables)
- Phi corrected (with indication of the magnitude of the effect size; only for 2x2 tables)
- Phi signed (with indication of the magnitude of the effect size; only for 2x2 tables)
- Adjusted contingency coefficient C (with indication of the magnitude of the effect)
- Maximum-corrected contingency coefficient C (with indication of the magnitude of the effect)
- Cramer's V (with 95perc confidence interval; includes indication of the magnitude of the effect)
- Cramer's V max (used to compute V corrected; for tables of any size)
- Cramer's V corrected (with indication of the magnitude of the effect)
- Cramer's V standardised (with indication of the magnitude of the effect)
- Bias-corrected Cramer's V (with indication of the magnitude of the effect)
- Cohen's w (with indication of the magnitude of the effect)
- W coefficient (includes 95perc confidence interval and magnitude of the effect)

Margin-free measures:

- Yule's Q (only for 2x2 tables, includes 95perc confidence interval, p-value, and indication of the magnitude of the effect)
- Yule's Y (only for 2x2 tables, includes 95perc confidence interval, p-value and indication of the magnitude of the effect)
- Odds ratio (for 2x2 tables, includes 95perc confidence interval, p value, and indication of the magnitude of the effect)
- Independent odds ratios (for tables larger than 2x2)

PRE (Proportional Reduction in Error) measures:

- Goodman-Kruskal's lambda (both asymmetric and symmetric, with 95perc confidence interval)
- Corrected version of lambda (both asymmetric and symmetric)
- Goodman-Kruskal's tau (asymmetric, with 95perc confidence interval)

Chi-square-based measures

C corrected

The package computes the (unadjusted) C coefficient as well as its adjusted version (C_{adj}), which is equal to C divided by $\sqrt{(k-1)/k}$, where k is the number of rows or columns, whichever is smaller (Sheskin 2011). The adjustment factor is equal to the maximum value C can attain on the basis of the table's size alone. However, that does not take into account the actual configuration of the marginals (Berry et al. 2018). Therefore, a 'proper' correction is based on dividing C by the value it achieves in the *chi-square-maximising table* (see above). This maximum-corrected version of C is computed by the package and reported as *C corrected* as opposed to *C adj*, which refers to the first above-described adjustment. When interpreting the magnitude of the association, (unadjusted) C is assessed against Cohen's thresholds adjusted by C's maximum possible value given the table's

marginals, while C adj and C corrected are assessed against the standard Cohen's thresholds (scaled by table's size; see below).

For more details on on maximum-corrected C see: Berry et al. 2018.

Phi corrected

To refine the phi coefficient, scholars have introduced a corrected version. It accounts for the fact that the original coefficient (1) does not always have a maximum achievable value of unity since it depends on the marginal configuration, and therefore (2) it is not directly comparable across tables with different marginals. To calculate phi-corrected, one first computes phi-max, which represents the maximum possible value of phi under the given marginal totals. Phi-corrected is equal to phi/phi-max.

For more details see: Cureton 1959; Liu 1980; Davenport et al. 1991; Rash et al. 2011; Alberti 2024.

See also *Chi-square-maximising table* above.

Cramer's V and its 95perc confidence interval

The calculation of the 95perc confidence interval around Cramer's V is based on Smithson 2003: 39-41, and builds on the R code made available by the author on the web (<http://www.michaelsmithson.online/stats/CIstuff/CI.htm>).

Cramér's V corrected

Cramér's V corrected addresses a fundamental limitation of the standard Cramér's V coefficient: its maximum achievable value depends on the marginal distribution of the table and rarely reaches 1.0, even when there is perfect association. This makes direct comparisons of V values across tables with different marginal configurations problematic.

To compute V corrected, the function first calculates V max, which represents the maximum possible value of Cramér's V that could be achieved given the observed marginal totals. V max is derived from the chi-square statistic computed on the *chi-square-maximising table* (see above), which represents the theoretical upper bound of association given the constraints imposed by the marginals. V corrected is then calculated as the ratio $V / V \text{ max}$.

This correction serves two important purposes:

- It scales V relative to its theoretically achievable maximum, allowing V corrected to range from 0 to 1 under any marginal configuration. A value of 1.0 indicates that the observed association is as strong as is possible given the marginal constraints.
- It enables meaningful comparisons of association strength across tables with different marginal distributions. Two tables may have identical uncorrected V values, but if one has more skewed marginals (and thus a lower V max), its V corrected will be higher, appropriately reflecting that the observed association represents a larger proportion of what is theoretically possible in that table.

The interpretation of V corrected follows standard Cohen thresholds scaled for table size (see below under *Indication of the magnitude of the association*), since the correction ensures the coefficient is already scaled relative to its maximum achievable value.

For theoretical background, see: Berry et al. 2018.

See also *Chi-square-maximising table* and *Phi corrected* above.

Cramer's V standardised

This version of the coefficient is computed on the standardised table (see above under *Table Standardisation*), which is returned and rendered by the function. Since a number of association measures, among which Cramer's V, are affected by the skewness in the marginal distributions, the original table is first standardised and then Cramer's V is computed.

The rationale of the use of standardised tables as basis to compute Cramer's V is that coefficients calculated on standardised tables are comparable across tables because the impact of different marginal distributions is controlled for.

The value obtained by subtracting the ratio Cramer's V to Cramer's V standardised from 1 gives an idea of the reduction of the magnitude of V due to the skewness of the marginal sums (multiplied by 100 can be interpreted as percentage reduction).

Bias-corrected Cramer's V

The bias-corrected Cramer's V is based on Bergsma 2013: 323–328.

W-hat coefficient

It addresses some limitations of Cramer's V. When the marginal probabilities are unevenly distributed, V may overstate the strength of the association, proving pretty high even when the overall association is weak. W is based on the distance between observed and expected frequencies. It uses the squared distance to adjust for the unevenness of the marginal distributions in the table. The indication of the magnitude of the association is based on Cohen 1988 (see above). Unlike Kvalseth 2018a, the calculation of the 95 percent confidence interval is based on a bootstrap approach, employing B resampled tables, and the 2.5th and 97.5th percentiles of the bootstrap distribution. B is set by default to 1000, but can be customised by the user.

For more details see: Kvalseth 2018a.

Indication of the magnitude of the association as indicated by the chi-squared-based coefficients

The function provides indication of the magnitude of the association (effect size) for the phi, phi signed, phi corrected, C, C adj, C corrected, Cramer's V, V corrected, V standardised, V bias-corrected, Cohen's w, and W-hat.

The verbal articulation of the effect size is based on Cohen 1988, with an enhanced approach (building on Olivier-Bell 2013) to account for the maximum achievable values of certain measures.

For all measures, the classification follows these principles:

1. For uncorrected measures (Phi, Phi signed, C, Cramer's V), Cohen's thresholds are adjusted (Olivier-Bell 2013) by multiplying them by the maximum achievable value of the coefficient given the marginal distribution of the table (Phi max, C max given table's marginals, V max). This accounts for the fact that, in these measures, the maximum achievable value is 1.0 only under specific marginal configurations.
2. For corrected measures (Phi corrected, C adj, C corrected, V corrected, V standardised), the standard Cohen's thresholds are used, as these measures are already scaled relative to their maximum values.

3. For measures using standard Cohen thresholds (V bias-corrected, W), the original thresholds scaled for table size are used without additional adjustment. These measures are not explicitly maximum-corrected but benefit from the table size scaling.

4. Cohen's w uses fixed thresholds (small 0.1, medium 0.3, large 0.5) regardless of table size.

For measures in groups 1-3 with input cross-tabs larger than 2x2, the Cohen's thresholds are scaled based on the table's df, which (as per Cohen 1988) correspond to the smaller between the rows and columns number, minus 1. On the basis of the table's df, the thresholds are calculated as follows:

small effect: $0.100 / \sqrt{\min(nr, nc) - 1}$

medium effect: $0.300 / \sqrt{\min(nr, nc) - 1}$

large effect: $0.500 / \sqrt{\min(nr, nc) - 1}$

where nr and nc are the number of rows and number of columns respectively, and $\min(nr, nc) - 1$ corresponds to the table's df. Essentially, the thresholds for a small, medium, and large effect are computed by dividing the Cohen's thresholds for a 2x2 table (df=1) by the square root of the input table's df.

This approach provides a nuanced interpretation of effect sizes that accounts for both table dimensions and marginal configurations. For example:

- An uncorrected V of 0.25 in a 3x4 table with V max = 0.6 would be compared to thresholds that are both scaled for the table's df (2) and adjusted by V max. The adjusted large threshold would be $(0.500/\sqrt{2}) \times 0.6 = 0.212$, making 0.25 a "large effect".

- A V corrected of 0.25 in the same table would be compared only to thresholds scaled for the table's df: 0.071 (small), 0.212 (medium), and 0.354 (large). Since 0.25 is between 0.212 and 0.354, it would be classified as a "medium effect".

- V bias-corrected and W would use the same thresholds as V corrected, without adjustment for maximum values.

The function also returns and render a detailed summary table of all thresholds used for interpreting effect sizes, providing transparency about the methodology and facilitating accurate reporting of results.

See: Cohen 1988; Sheskin 2011; Olivier and Bell 2013; Alberti 2024.

Margin-free measures

Yule's Q and Yule's Y

Yule's Q has a probabilistic interpretation. It tells us how much more likely is to draw pairs of individuals who share the same characteristics (concordant pairs) as opposed to drawing pairs who do not (discordant pairs).

Yule's Y represents the difference between the probabilities in the diagonal and off-diagonal cells, in the equivalent symmetric tables with equal marginals. In other words, Y measures the extent to which the probability of falling in the diagonal cells exceeds the probability of falling in the off-diagonal cells, in the standardised version of the table.

Note that, if either cross-tab's diagonal features 0, the *Haldane-Anscombe correction* is applied to both Q and Y. This results in a coefficient approaching, but not reaching, 1 in such circumstances, with the exact value depending on the table's frequencies. For more info on the correction, see the *Odds Ratio* section below.

On Yule's Q and Y, see Yule 1912.

On Yule's Y, see also Reynolds 1977.

On the probabilistic interpretation of Q, see: Davis 1971; Goodman-Kruskal 1979.

On the sensitivity of Yule's Q and Pearson's phi to different types of associations, see Alberti 2024.

Magnitude of the association as indicated by Yule's Q and Yule's Y

Given the relationship between Q and Y and the odds ratio, the magnitude of the association indicated by Q and Y is based on the thresholds proposed by Ferguson 2009 for the odds ratio (see below). Specifically, the thresholds for Q (in terms of absolute value) are:

- $|Q| < 0.330$ - Negligible
- $0.330 \leq |Q| < 0.500$ - Small
- $0.500 \leq |Q| < 0.600$ - Medium
- $|Q| \geq 0.600$ - Large

For Yule's Y (absolute value):

- $|Y| < 0.171$ - Negligible
- $0.171 \leq |Y| < 0.268$ - Small
- $0.268 \leq |Y| < 0.333$ - Medium
- $|Y| \geq 0.333$ - Large

As noted earlier, the effect size magnitude for Yule's Q and Yule's Y is based on the above odds ratio's thresholds.

Odds Ratio

For 2x2 tables, a single odds ratio is computed.

For tables larger than 2x2, independent odds ratios are computed for adjacent rows and columns, so representing the minimal set of ORs needed to describe all pairwise relationships in the contingency table.

For tables of size 2xk (where $k \geq 2$), pairwise odds ratios can be plotted (along with their confidence interval) by setting the `plor` parameter to TRUE (see the *Odd Ratios plot* section further down).

In all three scenarios, the *Haldane-Anscombe correction* is applied when zeros are present along any of the table's diagonals. This correction consists of adding 0.5 to every cell of the relevant 2x2 subtable before calculating the odds ratio. This ensures:

- For 2x2 tables: The correction is applied to the entire table if either diagonal contains a zero.
- For larger tables: The correction is applied to each 2x2 subtable used to calculate an independent odds ratio, if that subtable has a zero in either diagonal.
- For 2xk tables: The correction is applied to each pairwise comparison that involves a zero in either diagonal.

Note that, the *Haldane-Anscombe correction* results in a finite (rather than infinite) odds ratio, with the exact value depending on the table's frequencies.

On the Haldane-Anscombe correction, see: Fleiss et al 2003.

Odds Ratio effect size magnitude

The magnitude of the association indicated by the odds ratio is based on the thresholds (and corresponding reciprocals) suggested by Ferguson 2009 (for other thresholds, see for instance Chen et al 2010):

- $OR < 2.0$ - Negligible
- $2.0 \leq OR < 3.0$ - Small
- $3.0 \leq OR < 4.0$ - Medium
- $OR \geq 4.0$ - Large

Odd Ratios plot

For $2 \times k$ table, where $k \geq 2$:

by setting the `plot.or` parameter to `TRUE`, a plot showing the odds ratios and their 95percent confidence interval will be rendered. The confidence level can be modified via the `or.alpha` parameter.

The odds ratios are calculated for the column levels, and one of them is to be selected by the user as a reference for comparison via the `reference.level` parameter (set to 1 by default). Also, the user may want to select the row category to which the calculation of the odds ratios makes reference (using the `row.level` parameter, which is set to 1 by default).

As mentioned earlier, if any of the 2×2 subtables on which the odds ratio is calculated features zeros along any of the diagonal, the *Haldane-Anscombe* correction is applied.

To better understand the rationale of plotting the odds ratios, consider the following example, which uses on the famous Titanic data:

Create a 2×3 contingency table:

```
mytable <- matrix(c(123, 158, 528, 200, 119, 181), nrow = 2, byrow = TRUE)
colnames(mytable) <- c("1st", "2nd", "3rd")
rownames(mytable) <- c("Died", "Survived")
```

Now, we perform the test and visualise the odds ratios:

```
chisquare(mytable, plot.or=TRUE, reference.level=1, row.level=1)
```

In the rendered plot, we can see the odds ratios and confidence intervals for the second and third column level (i.e., 2nd class and 3rd class) because the first column level has been selected as reference level. The odds ratios are calculated making reference to the first row category (i.e., *Died*). From the plot, we can see that, compared to the 1st class, passengers on the 2nd class have 2.16 times larger odds of dying; passengers on the 3rd class have 4.74 times larger odds of dying compared to the 1st class.

Note that if we set the `row.level` parameter to 2, we make reference to the second row category, i.e. *Survived*:

```
chisquare(mytable, plot.or=TRUE, reference.level=1, row.level=2)
```

In the plot, we can see that passengers in the 2nd class have 0.46 times the odds of surviving of passengers in the 1st class, while passengers from the 3rd class have 0.21 times the odds of surviving of those travelling in the 1st class.

PRE (Proportional Reduction in Error) measures

Confidence Interval around Goodman-Kruskal's Lambda

Confidence Interval Calculation

The 95 percent confidence interval around lambda uses the asymptotic variance formula (see, e.g., Reynolds 1977:50, equation 2.32) and is *not* truncated at zero, even though lambda itself ranges

from 0 to 1. This approach correctly reflects sampling uncertainty. The variance formula assumes **multinomial sampling** (simple random sample). For tables with fixed marginals (e.g. experimental designs with fixed row or column totals), a different variance formula is needed (see Goodman-Kruskal 1972 for details).

Special Properties of Lambda's Sampling Distribution

Lambda has a degenerate sampling distribution at its boundaries. Specifically:

- If the population lambda = 0, then the sample lambda will **always** equal exactly 0
- If the population lambda = 1, then the sample lambda will **always** equal exactly 1

Inference Rules

Because of this degeneracy, hypothesis testing operates differently from standard measures:

1. **If sample lambda = 0 exactly:** We cannot reject the hypothesis that the population lambda = 0, regardless of confidence interval width. The observed value is consistent with no association.
2. **If sample lambda different from 0 (even by a tiny amount):** We can definitively rule out that the population lambda equals exactly zero, because a true zero would have produced a sample estimate of exactly 0. However, confidence intervals extending below zero indicate the population lambda might be **very close to** (but not equal to) zero.

Practical Interpretation Example

Suppose lambda = 0.033 with 95 percent CI [-0.004, 0.060]:

- **Conclusion:** The population lambda is not exactly zero (because we observed lambda = 0.033), but it could be very small (close to zero).
- **Substantive meaning:** There is statistically detectable predictive association, but it may be trivially small in magnitude.
- **Recommendation:** Always report both the point estimate and the confidence interval. Focus on the lower confidence interval bound to assess the minimum plausible association strength.

When Confidence Intervals Include Zero

A confidence interval that includes zero does *not* mean "no association" in the usual sense. Instead, it indicates:

- The population lambda is definitely not exactly zero (if sample lambda different from 0)
- But the population lambda could be negligibly small (close to zero)
- The width of the confidence interval below zero indicates imprecision in estimating a small non-zero parameter

For symmetric lambda and tau, similar logic applies, though their distributions are less degenerate than asymmetric lambda.

For more details see: Reynolds 1977:48-51; Goodman-Kruskal 1972; Bishop et al. 2007:388-392.

Corrected Goodman-Kruskal's lambda

The corrected Goodman-Kruskal's lambda adeptly addresses skewed or unbalanced marginal probabilities which create problems to the traditional lambda. By emphasizing categories with higher

probabilities through a process of squaring maximum probabilities and normalizing with marginal probabilities, this refined coefficient addresses inherent limitations of lambda.

For more details see: Kvalseth 2018b.

ADDITIONAL NOTES ON CALCULATIONS

Additional notes on calculations:

- the **Phi** coefficient is based on the chi-square statistic as per Sheskin 2011's equation 16.21, whereas the **Phi signed** is after Sheskin's equation 16.20;
- the **2-sided p value of Yule's Q** is calculated following Sheskin 2011's equation 16.24;
- **Cohen's w** is calculated as $V * \sqrt{\min(nr, nc) - 1}$, where V is Cramer's V , and nr and nc are the number of rows and columns respectively; see Sheskin 2011: 679;
- the **95perc confidence interval** around **Goodman-Kruskal's lambda** is calculated as per Reynolds 1977: 50;
- the **symmetric** version of **Goodman-Kruskal's lambda** is calculated as per Reynolds 1984: 55-57;
- **Goodman-Kruskal's tau** is calculated as per Reynolds 1984: 57-60;
- the **95perc confidence interval** around **Goodman-Kruskal's tau** is calculated as per Bishop et al. 2007: 391-392.

Value

The function produces **optional charts** (distribution of the permuted chi-square statistic and a plot of the odds ratios between a reference column level and the other ones, the latter only for 2xk tables where $k \geq 2$), and a number of **output tables** that are nicely formatted with the help of the *gt* package. Note that these tables are all automatically rendered when the `render.all.tables` parameter is set to TRUE (default is FALSE). The tables are returned in a list (see further down) so that they can be separately and/or individually rendered.

The output 'gt' tables are listed below:

- Input contingency table (with some essential analytical results annotated at the bottom)
- Expected frequencies
- Cells' chi-square value
- Cells' relative contribution (in percent) to the chi-square statistic (cells in RED feature a larger-than-average contribution)
- Cells' absolute contribution (in percent) to the chi-square statistic (colour same as above)
- Chi-square-maximising table (with indication of the associated chi-square value, that is, the maximum value of the chi-square statistic achievable given the table margins)
- Standardised residuals (RED for large significant residuals, BLUE for small significant residuals)
- Moment-corrected standardised residuals (colour same as above)
- Adjusted standardised residuals (colour same as above)
- Standardised median polish residuals

- Adjusted standardised median polish residuals
- Quetelet Indices
- IJ association factors
- PEM (Percentage of Maximum Deviation from Independence)
- Adjusted standardised counts
- Goodman-Kruskal residuals (column variable as independent)
- Goodman-Kruskal residuals (row variable as independent)
- Input contingency table standardised via Iterative Proportional Fitting
- Table of independent odds ratios (for tables larger than 2x2)
- Effect Size Interpretation Thresholds (Chi-Square-Based Measures)
- Table of output statistics, p values, and association measures

Also, the function returns a **list containing the following elements**:

- **input.table**:
 - *crosstab*: input contingency table.
- **chi.sq.maxim.table**:
 - *chi.sq.maximising.table*: chi-square-maximising table.
- **standardised.table**:
 - *standard.table*: standardised table on which Cramer's V standardised is computed.
- **chi.sq.related.results**:
 - *exp.freq*: table of expected frequencies.
 - *smallest.exp.freq*: smallest expected frequency.
 - *avrg.exp.freq*: average expected frequency.
 - *chisq.values*: cells' chi-square value.
 - *chisq.relat.contrib*: cells' relative contribution (in percent) to the chi-square statistic.
 - *chisq.abs.contrib*: cells' absolute contribution (in percent) to the chi-square statistic.
 - *chisq.statistic*: observed chi-square value.
 - *chisq.p.value*: p value of the chi-square statistic.
 - *chisq.max*: chi-square value computed on the chi-square-maximising table.
 - *chi.sq.power*: retrospective power of the traditional chi-square test.
 - *chisq.adj*: chi-square statistic adjusted using the (N-1)/N correction.
 - *chisq.adj.p.value*: p value of the adjusted chi-square statistic.
 - *chisq.p.value.perm*: permutation-based p value, based on B permuted tables (see Details).
 - *chisq.p.value.MC*: Monte Carlo p value, based on B random tables (see Details).
- **G.square**:
 - *Gsq.statistic*: observed G-square value.
 - *Gsq.p.value*: p value of the G-square statistic.
- **post.hoc**:
 - *stand.resid*: table of chi-square standardised residuals.

- *mom.corr.stand.resid*: table of moment-corrected standardised residuals.
- *adj.stand.resid*: table of adjusted standardised residuals.
- *stand.med.pol.resid*: table of standardised median polish residuals.
- *adj.stand.med.pol.resid*: table of adjusted standardised median polish residuals.
- *Quetelet.Index*: table of Quetelet indices.
- *IJ.assoc.fact.*: table of IJ association factors.
- *PEM*: table of PEM values plus bootstrap 95perc CI.
- *adj.stand.counts*: table of adjusted standardised counts.
- *GK.residuals.col*: table of Goodman-Kruskal residuals (column variable as independent)
- *GK.residuals.row*: table of Goodman-Kruskal residuals (row variable as independent)
- **margin.free.assoc.measures:**
 - *Yule's Q*: Q coefficient (only for 2x2 tables).
 - *Yule's Q CI lower boundary*: lower boundary of the 95perc CI.
 - *Yule's Q CI upper boundary*: upper boundary of the 95perc CI.
 - *Yule's Q p.value*: 2-tailed p value of Yule's Q.
 - *Yule's Y*: Y coefficient (only for 2x2 tables).
 - *Yule's Y CI lower boundary*: lower boundary of the 95perc CI.
 - *Yule's Y CI upper boundary*: upper boundary of the 95perc CI.
 - *Yule's Y p.value*: 2-tailed p value of Yule's Y.
 - *Odds ratio*: odds ratio (for 2x2 tables).
 - *Odds ratio CI lower boundary*: lower boundary of the 95perc CI.
 - *Odds ratio CI upper boundary*: upper boundary of the 95perc CI.
 - *Odds ratio p.value*: p value of the odds ratio.
 - *ORs*: table of independent odds ratios (for tables larger than 2x2).
- **chi.sq.based.assoc.measures:**
 - *Phi.signed*: signed Phi coefficient (only for 2x2 tables).
 - *Phi*: Phi coefficient (only for 2x2 tables).
 - *Phi.max*: maximum value of Phi given the marginals (only for 2x2 tables).
 - *Phi.corr*: maximum-corrected Phi coefficient (equal to Phi/Phi max; only for 2x2 tables).
 - *C*: contingency coefficient.
 - *C max given table's size*: maximum value of C given the size of the table.
 - *C.adj*: adjusted contingency coefficient C (equal to C/Cmax given table's size).
 - *C.max given table's marginals*: maximum value of C given the marginals.
 - *C.corr*: maximum-corrected C coefficient (equal to C/Cmax given table's marginals)
 - *Cramer's V*: Cramer's V coefficient.
 - *Cramer's V CI lower boundary*: lower boundary of the 95perc CI.
 - *Cramer's V CI upper boundary*: upper boundary of the 95perc CI.
 - *Cramer's V max*: Maximum value of Cramer's V given the marginals.
 - *Cramer's V corr*: corrected V coefficient (equal to V/Vmax).
 - *Cramer's V standard.*: Cramer's V computed on the standardised table.
 - *1-(Cramer's V/V standard.)*: value indicating the reduction of the magnitude of V due to the skewness of the marginal sums.

- *Cramer's Vbc*: bias-corrected Cramer's V coefficient.
- *w*: Cohen's w.
- *W*: W coefficient.
- *W CI lower boundary*: lower boundary of the bootstrap 95perc CI.
- *W CI upper boundary*: upper boundary of the bootstrap 95perc CI.

- **PRE.assoc.measures:**

- *lambda (rows dep.)*: Goodman-Kruskal's lambda coefficient (considering the rows being the dependent variable).
- *lambda (rows dep.) CI lower boundary*: lower boundary of the 95perc CI.
- *lambda (rows dep.) CI upper boundary*: upper boundary of the 95perc CI.
- *lambda (cols dep.)*: Goodman-Kruskal's lambda coefficient (considering the columns being the dependent variable).
- *lambda (cols dep.) CI lower boundary*: lower boundary of the 95perc CI.
- *lambda (cols dep.) CI upper boundary*: upper boundary of the 95perc CI.
- *lambda (symmetric)*: Goodman-Kruskal's symmetric lambda coefficient.
- *lambda (symmetric) CI lower boundary*: lower boundary of the 95perc CI.
- *lambda (symmetric) CI upper boundary*: upper boundary of the 95perc CI.
- *lambda corrected (rows dep.)*: corrected version of the lambda coefficient (considering the rows being the dependent variable).
- *lambda corrected (cols dep.)*: corrected version of the lambda coefficient (considering the columns being the dependent variable).
- *lambda corrected (symmetric)*: corrected version of the symmetric lambda coefficient.
- *tau (rows dep.)*: Goodman-Kruskal's tau coefficient (considering the rows being the dependent variable).
- *tau (rows dep.) CI lower boundary*: lower boundary of the 95perc CI.
- *tau (rows dep.) CI upper boundary*: upper boundary of the 95perc CI.
- *tau (cols dep.)*: Goodman-Kruskal's tau coefficient (considering the columns being the dependent variable).
- *tau (cols dep.) CI lower boundary*: lower boundary of the 95perc CI.
- *tau (cols dep.) CI upper boundary*: upper boundary of the 95perc CI.

- **gt.tables:**

- *input.table*: gt table of input contingency table
- *expected.frequencies*: gt table of expected frequencies
- *chisq.values*: gt table of cells' chi-square values
- *relative.contributions*: gt table of cells' relative contribution to chi-square
- *absolute.contributions*: gt table of cells' absolute contribution to chi-square
- *chisq.max.table*: gt table of chi-square-maximising table
- *standardised.residuals*: gt table of standardised residuals
- *moment.corrected.residuals*: gt table of moment-corrected standardised residuals
- *adjusted.standardised.residuals*: gt table of adjusted standardised residuals
- *stand.med.pol.residuals*: gt table of standardised median polish residuals
- *adj.stand.med.pol.residuals*: gt table of adjusted standardised median polish residuals

- *quetelet.index*: gt table of Quetelet indices
- *ij.association*: gt table of IJ association factors
- *PEM*: gt table of PEM values and bootstrap 95perc CI
- *adjusted.stand.counts*: gt table of adjusted standardised counts
- *gk.residuals.col*: gt table of Goodman-Kruskal residuals (column variable as dependent)
- *gk.residuals.row*: gt table of Goodman-Kruskal residuals (row variable as dependent)
- *standardised.table*: gt table of standardised input contingency table
- *odds.ratios*: gt table of independent odds ratios (only for tables larger than 2x2)
- *effect.size.thresholds*: gt table of effect size interpretation thresholds (for chi-square-based measures)
- *analysis.report*: gt table of output statistics

These gt tables can be accessed and re-rendered at any time using the standard print method (e.g., `print(results$gt.tables$input.table)`). They can also be exported to various formats using gt's `gtsave` function. For example:

- HTML: `gtsave(results$gt.tables$input.table, "mytable.html")`
- PDF/PNG: `gtsave(results$gt.tables$input.table, "mytable.pdf")` or `"mytable.png"` (requires the 'webshot2' package)
- Word: `gtsave(results$gt.tables$input.table, "mytable.docx")` (requires the 'rmarkdown' package)
- RTF: `gtsave(results$gt.tables$input.table, "mytable.rtf")`
- LaTeX: `gtsave(results$gt.tables$input.table, "mytable.tex")`

Note that the *p-values* returned in the above list are expressed in scientific notation, whereas the ones reported in the output table featuring the tests' result and measures of association are reported as broken down into classes (e.g., <0.05, or <0.01, etc), with the exception of the Monte Carlo p-value.

The **following examples**, which use in-built datasets, can be run to familiarise with the function:

```
-perform the test on the in-built 'social_class' dataset:
result <- chisquare(social_class)
```

```
-perform the test on a 2x2 subset of the 'diseases' dataset:
mytable <- diseases[3:4, 1:2]
result <- chisquare(mytable)
```

```
-perform the test on a 2x2 subset of the 'safety' dataset:
mytable <- safety[c(4,1), c(1,6)]
result <- chisquare(mytable)
```

```
-build a toy dataset in 'long' format (gender vs. opinion about death sentence):
mytable <- data.frame(GENDER=c(rep("F", 360), rep("M", 340)), OPINION=c(rep("oppose",
235), rep("favour", 125), rep("oppose", 160), rep("favour", 180)))

-perform the test specifying that the input table is in 'long' format:
result <- chisquare(mytable, format="long")
```

References

- Agresti, A. (2013). *Categorical Data Analysis* (3rd ed.). Wiley. ISBN 9780470463635.
- Agresti, A., Franklin, C., & Klingenberg, B. (2022). *Statistics: The Art and Science of Learning from Data*, (5th ed.). Pearson Education.
- Alberti, G. (2024). *From Data to Insights: A Beginner's Guide to Cross-Tabulation Analysis*. Routledge - CRC Press.
- Beh E.J., Lombardo R. 2014. *Correspondence Analysis: Theory, Practice and New Strategies*, Chichester, Wiley.
- Beh, E. J., & Lombardo, R. (2021). *An introduction to correspondence analysis*. John Wiley & Sons.
- Beasley TM and Schumacker RE. 1995. Multiple Regression Approach to Analyzing Contingency Tables: Post Hoc and Planned Comparison Procedures. *The Journal of Experimental Education*, 64(1).
- Bergsma, W. 2013. A bias correction for Cramér's V and Tschuprow's T. *Journal of the Korean Statistical Society*. 42 (3).
- Berry, K. J., Johnston, J. E., & Mielke, P. W., Jr. (2018). *The Measurement of Association: A Permutation Statistical Approach*. Springer.
- Bishop, Y. M., Fienberg, S. E., & Holland, P. W. (2007). *Discrete Multivariate Analysis: Theory and Practice*. Springer. ISBN 9780387728056
- Campbell, I. (2007). Chi-squared and Fisher–Irwin tests of two-by-two tables with small sample recommendations. In *Statistics in Medicine* (Vol. 26, Issue 19, pp. 3661–3675).
- Chen, H., Cohen, P., and Chen, S. (2010). How Big is a Big Odds Ratio? Interpreting the Magnitudes of Odds Ratios in Epidemiological Studies. In *Communications in Statistics - Simulation and Computation* (Vol. 39, Issue 4, pp. 860–864).
- Cibois, P. (1993). Le PEM, pourcentage de l'écart maximum: Un indice de liaison entre modalités d'un tableau de contingence. *Bulletin de Methodologie Sociologique*, 40, 43-63.
- Cohen, J. 1988. *Statistical power analysis for the behavioral sciences* (2nd ed). Hillsdale, N.J: L. Erlbaum Associates.
- Cureton, E. E. (1959). Note on phi/phimax. In *Psychometrika* (Vol. 24, Issue 1, pp. 89–91).
- Davenport, E. C., Jr., & El-Sanhury, N. A. (1991). Phi/Phimax: Review and Synthesis. In *Educational and Psychological Measurement* (Vol. 51, Issue 4, pp. 821–828).
- Davis, J. A. (1971). *Elementary Survey Analysis*. Prentice Hall. ISBN 9780132605472.
- Fagerland, M. W., Lydersen, S., & Laake, P. (2017). *Statistical Analysis of Contingency Tables*. CRC Press. ISBN 9781466588172.
- Ferguson, C. J. (2009). An effect size primer: A guide for clinicians and researchers. *Professional Psychology: Research and Practice*, 40(5), 532–538.
- Fienberg, S. E. (1971). A statistical technique for historians: Standardizing tables of counts. *The Journal of Interdisciplinary History*, 1(2), 305-315.
- Fleiss, J. L., Levin, B., & Paik, M. C. 2003. *Statistical Methods for Rates and Proportions* (3rd ed.). Wiley.
- Garcia-Perez, MA, and Nunez-Anton, V. 2003. Cellwise Residual Analysis in Two-Way Contingency Tables. *Educational and Psychological Measurement*, 63(5).

- Goodman, L. A. (1996). A Single General Method for the Analysis of Cross-Classified Data: Reconciliation and Synthesis of Some Methods of Pearson, Yule, and Fisher, and also Some Methods of Correspondence Analysis and Association Analysis. *Journal of the American Statistical Association*, 91(433), 408-428.
- Goodman, L. A., & Kruskal, W. H. (1972). Measures of association for cross classifications, IV: Simplification of asymptotic variances. *Journal of the American Statistical Association*, 67(338), 415-421.
- Goodman, L. A., & Kruskal, W. H. (1979). *Measures of Association for Cross Classifications*. Springer-Verlag. ISBN 9780387904436.
- Greenwood, P. E., & Nikulin, M. S. (1996). *A guide to chi-squared testing*. John Wiley & Sons.
- Haberman, S. J. (1973). The Analysis of Residuals in Cross-Classified Tables. In *Biometrics* (Vol. 29, Issue 1, p. 205).
- Kroonenberg, P. M., & Lombardo, R. (1999). Nonsymmetric correspondence analysis: A tool for analysing contingency tables with a dependence structure. *Multivariate Behavioral Research*, 34(3), 367-396.
- Kvålseth, T. O. (2018a). An alternative to Cramér's coefficient of association. In *Communications in Statistics - Theory and Methods* (Vol. 47, Issue 23, pp. 5662-5674).
- Kvålseth, T. O. (2018b). Measuring association between nominal categorical variables: an alternative to the Goodman-Kruskal lambda. In *Journal of Applied Statistics* (Vol. 45, Issue 6, pp. 1118-1132).
- Lefèvre, B., & Champely, S. (2009). Méthodes statistiques globales et locales d'analyse d'un tableau de contingence par les tailles d'effet et leurs intervalles de confiance. *Bulletin de Methodologie Sociologique*, 103, 50-65.
- Liu, R (1980). A Note on Phi-Coefficient Comparison. In *Research in Higher Education* (Vol. 13, No. 1, pp. 3-8).
- Mirkin, B. (2023). A straightforward approach to chi-squared analysis of associations in contingency tables. In E. J. Beh, R. Lombardo, & J. G. Clavel (Eds.), *Analysis of Categorical Data from Historical Perspectives (Behaviormetrics: Quantitative Approaches to Human Behavior, vol. 17)*. Springer.
- Mosteller, F., & Parunak, A. (1985). Identifying extreme cells in a sizable contingency table: Probabilistic and exploratory approaches. In D. C. Hoaglin, F. Mosteller, & J. W. Tukey (Eds.), *Exploring Data Tables, Trends, and Shapes* (pp. 189-224). New York: Wiley.
- Olivier, J., & Bell, M. L. (2013). Effect sizes for 2x2 contingency tables. *PLoS ONE*, 8(3), e58777.
- Oyeyemi, G. M., Adewara, A. A., Adebola, F. B., & Salau, S. I. (2010). On the Estimation of Power and Sample Size in Test of Independence. In *Asian Journal of Mathematics and Statistics* (Vol. 3, Issue 3, pp. 139-146).
- Phipson, B., & Smyth, G. K. (2010). Permutation p-values should never be zero: Calculating exact p-values when permutations are randomly drawn. *Statistical Applications in Genetics and Molecular Biology*, 9(1), Article 39.
- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in Psychology Using R and SPSS*. Wiley.
- Reynolds, H. T. 1977. *The Analysis of Cross-Classifications*. New York: Free Press.

- Reynolds, H. T. 1984. Analysis of Nominal Data (Quantitative Applications in the Social Sciences) (1st ed.). SAGE Publications.
- Rhoades, H. M., & Overall, J. E. (1982). A sample size correction for Pearson chi-square in 2x2 contingency tables. In Psychological Bulletin (Vol. 91, Issue 2, pp. 418–423).
- Richardson, J. T. E. (2011). The analysis of 2×2 contingency tables-Yet again. In Statistics in Medicine (Vol. 30, Issue 8, pp. 890–890).
- Rosenthal, R., & Rosnow, R. L. (2008). Essentials of Behavioral Research: Methods and Data Analysis (3rd ed.). McGraw-Hill Higher Education.
- Roscoe, J. T., & Byars, J. A. (1971). An Investigation of the Restraints with Respect to Sample Size Commonly Imposed on the Use of the Chi-Square Statistic. Journal of the American Statistical Association, 66(336), 755–759.
- Sheskin, D. J. 2011. Handbook of Parametric and Nonparametric Statistical Procedures, Fifth Edition (5th ed.). Chapman and Hall/CRC.
- Simonoff, J. S. (2003). Analyzing Categorical Data. New York: Springer.
- Smith, K. W. (1976). Marginal Standardization and Table Shrinking: Aids in the Traditional Analysis of Contingency Tables. Social Forces, 54(3), 669-693.
- Smithson M.J. 2003. Confidence Intervals, Quantitative Applications in the Social Sciences Series, No. 140. Thousand Oaks, CA: Sage.
- Upton, G. J. G. (1982). A Comparison of Alternative Tests for the 2×2 Comparative Trial. In Journal of the Royal Statistical Society. Series A (General) (Vol. 145, Issue 1, p. 86).
- Yule, G. U. (1912). On the methods of measuring association between two attributes. Journal of the Royal Statistical Society, 75(6), 579–652.
- Zar, J. H. (2014). Biostatistical analysis (5th ed.). Pearson New International Edition.

Examples

```
# Perform the analysis on a 2x2 subset of the in-built 'social_class' dataset
result <- chisquare(social_class[c(1:2), c(1:2)], B=9)
```

diseases

Dataset: Cross-tabulation of quantity of tobacco smoked daily vs. cause of death

Description

Cross-tabulation (15x4) of the amount of tobacco smoked on a daily basis (in grams) against cause of death.

After: Velleman P F, Hoaglin D C, Applications, Basics, and Computing of Exploratory Data Analysis, Wadsworth Pub Co 1984 (Exhibit 8-1)

Usage

```
data(diseases)
```

Format

dataframe

safety

Dataset: Cross-tabulation of people's feeling of safety vs. town size

Description

Cross-tabulation (4x6).

Usage

`data(safety)`

Format

dataframe

social_class

Dataset: Cross-tabulation of social class vs. diagnostic category for a sample of psychiatric patients

Description

Cross-tabulation (3x4) after: Everitt B.S (1992), The Analysis of Contingency Tables, Chapman&Hall/CRC, second edition, table 3.13.

Usage

`data(social_class)`

Format

dataframe

Index

- * **chiperm**

- chisquare, [2](#)

- * **datasets**

- diseases, [26](#)

- safety, [27](#)

- social_class, [27](#)

chisquare, [2](#)

diseases, [26](#)

gtsave, [23](#)

safety, [27](#)

social_class, [27](#)