

coin: A Computational Framework for Conditional Inference

Torsten Hothorn¹, Kurt Hornik², Mark van de Wiel³
and Achim Zeileis²

¹Institut für Medizininformatik, Biometrie und Epidemiologie
Friedrich-Alexander-Universität Erlangen-Nürnberg
Waldstraße 6, D-91054 Erlangen, Germany
`Torsten.Hothorn@R-project.org`

²Department für Statistik und Mathematik, Wirtschaftsuniversität Wien
Augasse 2-6, A-1090 Wien, Austria
`Kurt.Hornik@R-project.org`
`Achim.Zeileis@R-project.org`

³Department of Mathematics, Vrije Universiteit
De Boelelaan 1081a, 1081 HV Amsterdam, The Netherlands
`mark.vdwiel@vumc.nl`

1 Introduction

The **coin** package provides a unified approach to conditional inference procedures commonly known as *permutation tests*. The theoretical basis for the design and implementation is the general framework for permutation tests given by [Strasser and Weber \(1999\)](#). For a very flexible formulation of multivariate linear statistics, [Strasser and Weber \(1999\)](#) derived the conditional expectation and covariance of the conditional (permutation) distribution as well as the multivariate limiting distribution. Test procedures for nominal, ordinal and continuous data with or without censoring are all part of this framework. For a more detailed overview see [Hothorn *et al.* \(2006\)](#).

The conceptual Strasser-Weber framework of permutation tests for arbitrary independence and symmetry problems are available via the generic functions `independence_test` and `symmetry_test` respectively. In addition to these very flexible procedures, a large set of convenience functions including well-known as well as lesser-known classical and non-classical procedures for tests of independence between two continuous variables, two- and K -sample tests for location and scale alternatives, tests of independence for contingency tables as well as tests of marginal homogeneity and symmetry have been implemented within the framework. Currently, the following conditional test procedures are available:

<code>spearman_test</code>	Spearman correlation test
<code>fisyat_test</code>	Fisher-Yates correlation test
<code>quadrant_test</code>	Quadrant test
<code>koziol_test</code>	Koziol-Nemec test
<code>oneway_test</code>	Two- and K -sample Fisher-Pitman permutation test
<code>wilcox_test</code>	Wilcoxon-Mann-Whitney test
<code>kruskal_test</code>	Kruskal-Wallis test
<code>normal_test</code>	Two- and K -sample van der Waerden test
<code>median_test</code>	Two- and K -sample Brown-Mood median test
<code>savage_test</code>	Two- and K -sample Savage test
<code>taha_test</code>	Two- and K -sample Taha test
<code>klotz_test</code>	Two- and K -sample Klotz test
<code>mood_test</code>	Two- and K -sample Mood test
<code>ansari_test</code>	Two- and K -sample Ansari-Bradley test
<code>fligner_test</code>	Two- and K -sample Fligner-Killeen test
<code>conover_test</code>	Two- and K -sample Conover-Iman test
<code>logrank_test</code>	Two- and K -sample weighted logrank tests (Logrank test, Gehan-Breslow test, Tarone-Ware test, ...)
<code>sign_test</code>	Sign test
<code>wilcoxsign_test</code>	Wilcoxon signed-rank test
<code>friedman_test</code>	Friedman test and Page test
<code>quade_test</code>	Quade test
<code>chisq_test</code>	Pearson χ^2 test
<code>cmh_test</code>	Generalized Cochran-Mantel-Haenszel test
<code>lbl_test</code>	Linear-by-linear association test
<code>mh_test</code>	Marginal homogeneity tests (McNemar test, Cochran Q test, Stuart-Maxwell test, ...)
<code>maxstat_test</code>	Generalized maximally selected statistics

These convenience functions essentially perform a certain transformation of the data, e.g., a rank transformation, and then call `independence_test` or `symmetry_test` for the computation of the (multivariate) linear statistic and its conditional expectation and covariance as well as the scalar test statistic and its null distribution. The exact null distribution can be approximated by the asymptotic distribution or via conditional Monte Carlo resampling for all test procedures, whereas the exact null distribution is available for special cases only. Moreover, all test procedures allow for the specification of blocks pertaining to, e.g., stratification or repeated measurements.

2 Permutation Tests

In the following we assume that we are provided with n observations

$$(\mathbf{Y}_i, \mathbf{X}_i, w_i, b_i), \quad i = 1, \dots, n.$$

The variables $\mathbf{Y}$ and $\mathbf{X}$ from sample spaces $\mathcal{Y}$ and $\mathcal{X}$ may be measured at arbitrary scales and may be multivariate as well. In addition to those measurements, case weights w and a factor b coding blocks may be available. For the sake of simplicity, we assume $w_i = 1$ and $b_i = 0$ for all observations $i = 1, \dots, n$ for the moment.

We are interested in testing the null hypothesis of independence of $\mathbf{Y}$ and $\mathbf{X}$

$$H_0 : D(\mathbf{Y}|\mathbf{X}) = D(\mathbf{Y})$$

against arbitrary alternatives. [Strasser and Weber \(1999\)](#) suggest to derive scalar test statistics for testing H_0 from multivariate linear statistics of the form

$$\mathbf{T} = \text{vec} \left(\sum_{i=1}^n w_i g(\mathbf{X}_i) h(\mathbf{Y}_i, (\mathbf{Y}_1, \dots, \mathbf{Y}_n))^\top \right) \in \mathbb{R}^{pq}. \quad (1)$$

Here, $g : \mathcal{X} \rightarrow \mathbb{R}^p$ is a transformation of $\mathbf{X}$ known as the *regression function* and $h : \mathcal{Y} \times \mathcal{Y}^n \rightarrow \mathbb{R}^q$ is a transformation of $\mathbf{Y}$ known as the *influence function*, where the latter may depend on $(\mathbf{Y}_1, \dots, \mathbf{Y}_n)$ in a permutation symmetric way. We will give specific examples on how to choose g and h later on.

The distribution of $\mathbf{T}$ depends on the joint distribution of $\mathbf{Y}$ and $\mathbf{X}$, which is unknown under almost all practical circumstances. At least under the null hypothesis one can dispose of this dependency by fixing $\mathbf{X}_1, \dots, \mathbf{X}_n$ and conditioning on all possible permutations S of the responses $\mathbf{Y}_1, \dots, \mathbf{Y}_n$. This principle leads to test procedures known as *permutation tests*.

The conditional expectation $\mu \in \mathbb{R}^{pq}$ and covariance $\Sigma \in \mathbb{R}^{pq \times pq}$ of $\mathbf{T}$ under H_0 given all permutations $\sigma \in S$ of the responses are derived by [Strasser and Weber \(1999\)](#):

$$\begin{aligned} \mu &= \mathbb{E}(\mathbf{T}|S) = \text{vec} \left(\left(\sum_{i=1}^n w_i g(\mathbf{X}_i) \right) \mathbb{E}(h|S)^\top \right), \\ \Sigma &= \mathbb{V}(\mathbf{T}|S) \\ &= \frac{\mathbf{w}_\cdot}{\mathbf{w}_\cdot - 1} \mathbb{V}(h|S) \otimes \left(\sum_i w_i g(\mathbf{X}_i) \otimes w_i g(\mathbf{X}_i)^\top \right) \\ &\quad - \frac{1}{\mathbf{w}_\cdot - 1} \mathbb{V}(h|S) \otimes \left(\sum_i w_i g(\mathbf{X}_i) \right) \otimes \left(\sum_i w_i g(\mathbf{X}_i) \right)^\top \end{aligned} \quad (2)$$

where $\mathbf{w}_\cdot = \sum_{i=1}^n w_i$ denotes the sum of the case weights, and $\otimes$ is the Kronecker product. The conditional expectation of the influence function is

$$\mathbb{E}(h|S) = \mathbf{w}_\cdot^{-1} \sum_i w_i h(\mathbf{Y}_i, (\mathbf{Y}_1, \dots, \mathbf{Y}_n)) \in \mathbb{R}^q$$

with corresponding $q \times q$ covariance matrix

$$\begin{aligned} \mathbb{V}(h|S) &= \mathbf{w}_\cdot^{-1} \sum_i w_i (h(\mathbf{Y}_i, (\mathbf{Y}_1, \dots, \mathbf{Y}_n)) - \mathbb{E}(h|S)) \\ &\quad (h(\mathbf{Y}_i, (\mathbf{Y}_1, \dots, \mathbf{Y}_n)) - \mathbb{E}(h|S))^\top. \end{aligned}$$

Having the conditional expectation and covariance at hand we are able to standardize a linear statistic $\mathbf{T} \in \mathbb{R}^{pq}$ of the form (1). Univariate test statistics c mapping an observed linear statistic $\mathbf{t} \in \mathbb{R}^{pq}$ into the real line can be of arbitrary form. An obvious choice is the maximum of the absolute values of the standardized linear statistic

$$c_{\max}(\mathbf{t}, \mu, \Sigma) = \max \left| \frac{\mathbf{t} - \mu}{\text{diag}(\Sigma)^{1/2}} \right|$$

utilizing the conditional expectation μ and covariance matrix Σ . The application of a quadratic form $c_{\text{quad}}(\mathbf{t}, \mu, \Sigma) = (\mathbf{t} - \mu) \Sigma^+ (\mathbf{t} - \mu)^\top$ is one alternative, although computationally more expensive because the Moore-Penrose inverse Σ^+ of Σ is involved.

The definition of one- and two-sided p -values used for the computations in the **coin** package is

$$\begin{aligned} P(c(\mathbf{T}, \mu, \Sigma) &\leq c(\mathbf{t}, \mu, \Sigma)) && \text{(less)} \\ P(c(\mathbf{T}, \mu, \Sigma) &\geq c(\mathbf{t}, \mu, \Sigma)) && \text{(greater)} \\ P(|c(\mathbf{T}, \mu, \Sigma)| &\geq |c(\mathbf{t}, \mu, \Sigma)|) && \text{(two-sided)}. \end{aligned}$$

Note that for quadratic forms only two-sided p -values are available and that in the one-sided case maximum type test statistics are replaced by

$$\min \left(\frac{\mathbf{t} - \mu}{\text{diag}(\Sigma)^{1/2}} \right) \quad \text{(less)} \quad \text{and} \quad \max \left(\frac{\mathbf{t} - \mu}{\text{diag}(\Sigma)^{1/2}} \right) \quad \text{(greater)}.$$

The conditional distribution and thus the p -value of the statistics $c(\mathbf{t}, \mu, \Sigma)$ can be computed in several different ways. For some special forms of the linear statistic, the exact distribution of the test statistic is tractable. For two-sample problems, the shift-algorithm by [Streitberg and Röhmel \(1986\)](#) and [Streitberg and Röhmel \(1987\)](#) and the split-up algorithm by [van de Wiel \(2001\)](#) are implemented as part of the package. Conditional Monte Carlo procedures can be used to approximate the exact distribution. [Strasser and Weber \(1999\)](#) proved (Theorem 2.3) that the conditional distribution of linear statistics $\mathbf{T}$ with conditional expectation μ and covariance Σ tends to a multivariate normal distribution with parameters μ and Σ as $n, \mathbf{w}. \rightarrow \infty$. Thus, the asymptotic conditional distribution of test statistics of the form $c_{\max}$ is normal and can be computed directly in the univariate case ($pq = 1$) or approximated by means of quasi-randomized Monte Carlo procedures in the multivariate setting ([Genz, 1992](#)). For quadratic forms c_{quad} which follow a χ^2 distribution with degrees of freedom given by the rank of Σ (see [Johnson and Kotz, 1970](#), Chapter 29), exact probabilities can be computed efficiently.

3 Illustrations and Applications

The main workhorse `independence_test` essentially allows for the specification of $\mathbf{Y}, \mathbf{X}$ and b through a formula interface of the form `y ~ x | b`, case weights can be defined by a formula with one variable on the right hand side only. Four additional arguments are available for the specification of the regression function g (`xtrans`), the influence function h (`ytrans`), the form of the test statistic c (`teststat`) and the null distribution (`distribution`).

Independent K -Sample Problems. When we want to compare the distribution of an univariate qualitative response $\mathbf{Y}$ in K groups given by a factor $\mathbf{X}$ at K levels, the transformation g is the dummy matrix coding the groups and h is either the identity transformation or a some form of rank transformation.

For example, the Kruskal-Wallis test may be computed as follows (example taken from [Hollander and Wolfe, 1999](#), Table 6.3, page 200):

```
R> library("coin")
R> YOY <- data.frame(
  length = c(46, 28, 46, 37, 32, 41, 42, 45, 38, 44,
             42, 60, 32, 42, 45, 58, 27, 51, 42, 52,
             38, 33, 26, 25, 28, 28, 26, 27, 27, 27,
             31, 30, 27, 29, 30, 25, 25, 24, 27, 30),
  site = gl(4, 10, labels = as.roman(1:4))
```

```

)
R> it <- independence_test(length ~ site, data = YOY,
  ytrafo = function(data)
    trafo(data, numeric_trafo = rank_trafo),
  teststat = "quadratic")
R> it

```

Asymptotic General Independence Test

```

data: length by site (I, II, III, IV)
chi-squared = 22.852, df = 3, p-value = 4.335e-05

```

The linear statistic $\mathbf{T}$ is the sum of the ranks in each group and can be extracted via

```
R> statistic(it, type = "linear")
```

```

I    278
II   307
III  119
IV   116

```

Note that `statistic(..., type = "linear")` currently returns the linear statistic in matrix form, i.e.

$$\sum_{i=1}^n w_i g(\mathbf{X}_i) h(\mathbf{Y}_i, (\mathbf{Y}_1, \dots, \mathbf{Y}_n))^{\top} \in \mathbb{R}^{p \times q}.$$

The conditional expectation and covariance are available from

```
R> expectation(it)
```

```

I    205
II   205
III  205
IV   205

```

```
R> covariance(it)
```

```

      I      II      III      IV
I  1019.0385 -339.6795 -339.6795 -339.6795
II -339.6795 1019.0385 -339.6795 -339.6795
III -339.6795 -339.6795 1019.0385 -339.6795
IV -339.6795 -339.6795 -339.6795 1019.0385

```

and the standardized linear statistic $(\mathbf{T} - \mu) \text{diag}(\Sigma)^{-1/2}$ is

```
R> statistic(it, type = "standardized")
```

```

I    2.286797
II   3.195250
III -2.694035
IV  -2.788013

```

Since a quadratic form of the test statistic was requested via `teststat = "quadratic"`, the test statistic is

```
R> statistic(it)
```

```
[1] 22.85242
```

By default, the asymptotic distribution of the test statistic is computed, the p -value is

```
R> pvalue(it)
```

```
[1] 4.334659e-05
```

Life is much simpler with convenience functions very similar to those available in package **stats** for a long time. Here, the exact null distribution of the Kruskal-Wallis test is approximated by Monte Carlo resampling using 10000 replicates via

```
R> kt <- kruskal_test(length ~ site, data = YOY,  
                     distribution = approximate(nresample = 10000))
```

```
R> kt
```

Approximative Kruskal-Wallis Test

```
data: length by site (I, II, III, IV)  
chi-squared = 22.852, p-value < 1e-04
```

with p -value (and 99 % confidence interval) of

```
R> pvalue(kt)
```

```
[1] <1e-04
```

```
99 percent confidence interval:  
0.0000000000 0.0005296914
```

Of course it is possible to choose a $c_{\max}$ type test statistic instead of a quadratic form.

Independence in Contingency Tables. Independence in general two- or three-dimensional contingency tables can be tested by the Cochran-Mantel-Haenszel test. Here, both g and h are dummy matrices (example data from [Agresti, 2002](#), Table 7.8, page 288):

```
R> data("jobsatisfaction", package = "coin")  
R> ct <- cmh_test(jobsatisfaction)  
R> ct
```

Asymptotic Generalized Cochran-Mantel-Haenszel Test

```
data: Job.Satisfaction by  
      Income (<5000, 5000-15000, 15000-25000, >25000)  
      stratified by Gender  
chi-squared = 10.2, df = 9, p-value = 0.3345
```

The standardized contingency table allowing for an inspection of the deviation from the null hypothesis of independence of income and jobsatisfaction (stratified by gender) is

```
R> statistic(ct, type = "standardized")
```

	Very Dissatisfied	A Little Satisfied
<5000	1.3112789	0.69201053
5000-15000	0.6481783	0.83462550
15000-25000	-1.0958361	-1.50130926
>25000	-1.0377629	-0.08983052

	Moderately Satisfied	Very Satisfied
<5000	-0.2478705	-0.9293458
5000-15000	0.5175755	-1.6257547
15000-25000	0.2361231	1.4614123
>25000	-0.5946119	1.2031648

Ordered Alternatives. Of course, both job satisfaction and income are ordered variables. When $\mathbf{Y}$ is measured at J levels and $\mathbf{X}$ at K levels, $\mathbf{Y}$ and $\mathbf{X}$ are associated with score vectors $\eta \in \mathbb{R}^J$ and $\gamma \in \mathbb{R}^K$, respectively. The linear statistic is now a linear combination of the linear statistic $\mathbf{T}$ of the form

$$\mathbf{MT} = \text{vec} \left(\sum_{i=1}^n w_i \gamma^\top g(\mathbf{X}_i) (\eta^\top h(\mathbf{Y}_i, (\mathbf{Y}_1, \dots, \mathbf{Y}_n))^\top \right) \in \mathbb{R} \text{ with } \mathbf{M} = \eta \otimes \gamma.$$

By default, scores are $\eta = 1, \dots, J$ and $\gamma = 1, \dots, K$.

```
R> lbl_test(jobsatisfaction)
```

Asymptotic Linear-by-Linear Association Test

```
data: Job.Satisfaction (ordered) by
      Income (<5000 < 5000-15000 < 15000-25000 < >25000)
      stratified by Gender
Z = 2.5736, p-value = 0.01006
alternative hypothesis: two.sided
```

The scores η and γ can be specified to the linear-by-linear association test via a list whose names correspond to the variable names

```
R> lbl_test(jobsatisfaction,
            scores = list(Job.Satisfaction = c(1, 3, 4, 5),
                          Income = c(3, 10, 20, 35)))
```

Asymptotic Linear-by-Linear Association Test

```
data: Job.Satisfaction (ordered) by
      Income (<5000 < 5000-15000 < 15000-25000 < >25000)
      stratified by Gender
Z = 2.4812, p-value = 0.01309
alternative hypothesis: two.sided
```

Alternatively, scores can be specified through `of_trafo`

```
R> lbl_test(jobsatisfaction,
  ytrafo = function(data)
    trafo(data, ordered_trafo = function(y)
      of_trafo(y, scores = c(1, 3, 4, 5))),
  xtrafo = function(data)
    trafo(data, ordered_trafo = function(x)
      of_trafo(x, scores = c(3, 10, 20, 35))))
```

Asymptotic Linear-by-Linear Association Test

```
data: Job.Satisfaction (ordered) by
      Income (<5000 < 5000-15000 < 15000-25000 < >25000)
      stratified by Gender
Z = 2.4812, p-value = 0.01309
alternative hypothesis: two.sided
```

Incomplete Randomised Blocks. [Rayner and Best \(2001\)](#), Chapter 7, discuss the application of Durbin's test to data from sensory experiments, where incomplete block designs are common. As an example, data from taste-testing on ten dried eggs where mean scores for off-flavour from seven judges are given and one wants to assess whether there is any difference in the scores between the ten egg samples. The sittings are a block variable which can be added to the formula via '|'.

```
R> egg_data <- data.frame(
  scores = c(9.7, 8.7, 5.4, 5.0, 9.6, 8.8, 5.6, 3.6, 9.0,
    7.3, 3.8, 4.3, 9.3, 8.7, 6.8, 3.8, 10.0, 7.5,
    4.2, 2.8, 9.6, 5.1, 4.6, 3.6, 9.8, 7.4, 4.4,
    3.8, 9.4, 6.3, 5.1, 2.0, 9.4, 9.3, 8.2, 3.3,
    8.7, 9.0, 6.0, 3.3, 9.7, 6.7, 6.6, 2.8, 9.3,
    8.1, 3.7, 2.6, 9.8, 7.3, 5.4, 4.0, 9.0, 8.3,
    4.8, 3.8, 9.3, 8.3, 6.3, 3.8),
  sitting = factor(rep(c(1:15), rep(4, 15))),
  product = factor(c(1, 2, 4, 5, 2, 3, 6, 10, 2, 4, 6, 7,
    1, 3, 5, 7, 1, 4, 8, 10, 2, 7, 8, 9,
    2, 5, 8, 10, 5, 7, 9, 10, 1, 2, 3, 9,
    4, 5, 6, 9, 1, 6, 7, 10, 3, 4, 9, 10,
    1, 6, 8, 9, 3, 4, 7, 8, 3, 5, 6, 8)))
)
R> independence_test(scores ~ product | sitting,
  data = egg_data, teststat = "quadratic",
  ytrafo = function(data)
    trafo(data, numeric_trafo = rank_trafo,
      block = egg_data$sitting))
```

Asymptotic General Independence Test

```
data: scores by
      product (1, 2, 3, 4, 5, 6, 7, 8, 9, 10)
      stratified by sitting
chi-squared = 39.12, df = 9, p-value = 1.096e-05
```

and the Monte Carlo p -value can be computed via

```
R> pvalue(independence_test(scores ~ product | sitting,
  data = egg_data, teststat = "quadratic",
  ytrafo = function(data)
    trafo(data, numeric_trafo = rank_trafo,
      block = egg_data$sitting),
  distribution = approximate(nresample = 19999)))

[1] <5.00025e-05
99 percent confidence interval:
 0.000000000 0.000264894
```

If we assume that the products are ordered, the Page test is appropriate and can be computed as follows

```
R> independence_test(scores ~ product | sitting,
  data = egg_data,
  ytrafo = function(data)
    trafo(data, numeric_trafo = rank_trafo,
      block = egg_data$sitting),
  scores = list(product = 1:10))

Asymptotic General Independence Test

data:  scores by
      product (1 < 2 < 3 < 4 < 5 < 6 < 7 < 8 < 9 < 10)
      stratified by sitting
Z = -6.2166, p-value = 5.081e-10
alternative hypothesis: two.sided
```

Multiple Tests. One may be interested in testing multiple hypotheses simultaneously, either by using a linear combination of the linear statistic **KT**, or by specifying multivariate variables **Y** and/or **X**. For example, all pair comparisons may be implemented via

```
R> it <- independence_test(length ~ site, data = YOY,
  xtrafo = mcp_trafo(site = "Tukey"),
  teststat = "maximum",
  distribution = "approximate")

R> pvalue(it)

[1] 1e-04
99 percent confidence interval:
 5.012541e-07 7.427741e-04

R> pvalue(it, method = "single-step")

II - I    0.6618
III - I   0.0399
IV - I    0.0228
III - II  0.0001
IV - II   0.0001
IV - III  0.9989
```

When either g or h are multivariate, single-step adjusted p -values based on maximum-type statistics are computed as described in Westfall and Young (1993), algorithm 2.5 (page 47) and equation (2.8), page 50. Note that for the example shown above only the *minimum* p -value is adjusted appropriately because the subset pivotality condition is violated, i.e., the distribution of the test statistics under the complete null-hypothesis of no treatment effect of `site` is the basis of all adjustments instead of the corresponding partial null-hypothesis.

4 Quality Assurance

The test procedures implemented in package `coin` are continuously checked against results obtained by the corresponding implementations in package `stats` (where available). In addition, the test statistics and exact, approximative and asymptotic p -values for data examples given in the **StatXact** -6 user manual (Cytel Inc., 2003) are compared with the results reported in the **StatXact** 6 manual. Step-down multiple adjusted p -values have been checked against results reported by `mt.maxT` from package `multtest` (Pollard *et al.*, 2008). For details on the test procedures we refer to the R transcript files in directory `coin/tests`.

5 Acknowledgements

We would like to thank Helmut Strasser for discussions on the theoretical framework. Henric Winell provided clarification and examples for the Stuart-Maxwell test.

References

- Agresti A (2002). *Categorical Data Analysis*. 2nd edition. John Wiley & Sons, Hoboken, New Jersey.
- Cytel Inc (2003). **StatXact 6: Statistical Software for Exact Nonparametric Inference**. Cytel Software Corporation, Cambridge, MA, U.S.A. URL <https://cytel.com/>.
- Genz A (1992). “Numerical Computation of Multivariate Normal Probabilities.” *Journal of Computational and Graphical Statistics*, **1**(2), 141–149. doi10.1080/10618600.1992.10477010.
- Hollander M, Wolfe DA (1999). *Nonparametric Statistical Inference*. 2nd edition. John Wiley & Sons, New York, U.S.A.
- Hothorn T, Hornik K, van de Wiel MA, Zeileis A (2006). “A Lego System for Conditional Inference.” *The American Statistician*, **60**(3), 257–263. doi10.1198/000313006X118430.
- Johnson NL, Kotz S (1970). *Distributions in Statistics: Continuous Univariate Distributions 2*. John Wiley & Sons, New York.
- Pollard KS, Ge Y, Dudoit S (2008). **multtest: Resampling-Based Multiple Hypothesis Testing**. R package version 1.21.1, URL <https://CRAN.R-project.org/package=multtest>.
- Rayner JCW, Best DJ (2001). *A Contingency Table Approach to Nonparametric Testing*. Chapman & Hall, New York, U.S.A.
- Strasser H, Weber C (1999). “On the Asymptotic Theory of Permutation Statistics.” *Mathematical Methods of Statistics*, **8**(2), 220–250.

- Streitberg B, Röhm J (1986). “Exact Distributions for Permutation and Rank Tests: An Introduction to Some Recently Published Algorithms.” *Statistical Software Newsletter*, **12**(1), 10–17. ISSN 1609-3631.
- Streitberg B, Röhm J (1987). “Exakte Verteilungen für Rang- und Randomisierungstests im allgemeinen c -Stichprobenfall.” *EDV in Medizin und Biologie*, **18**(1), 12–19.
- van de Wiel MA (2001). “The Split-Up Algorithm: A Fast Symbolic Method for Computing p -Values of Distribution-Free Statistics.” *Computational Statistics*, **16**, 519–538.
- Westfall PH, Young SS (1993). *Resampling-Based Multiple Testing*. John Wiley & Sons, New York, U.S.A.