

Package ‘misty’

September 6, 2025

Type Package

Title Miscellaneous Functions 'T. Yanagida'

Version 0.7.5

Date 2025-09-06

Author Takuya Yanagida [aut, cre]

Maintainer Takuya Yanagida <takuya.yanagida@univie.ac.at>

Description Miscellaneous functions for (1) data handling (e.g., grand-mean and group-mean centering, coding variables and reverse coding items, scale and cluster scores, reading and writing Excel and SPSS files), (2) descriptive statistics (e.g., frequency table, cross tabulation, effect size measures), (3) missing data (e.g., descriptive statistics for missing data, missing data pattern, Little's test of Missing Completely at Random, and auxiliary variable analysis), (4) multilevel data (e.g., multilevel descriptive statistics, within-group and between-group correlation matrix, multilevel confirmatory factor analysis, level-specific fit indices, cross-level measurement equivalence evaluation, multilevel composite reliability, and multilevel R-squared measures), (5) item analysis (e.g., confirmatory factor analysis, coefficient alpha and omega, between-group and longitudinal measurement equivalence evaluation), (6) statistical analysis (e.g., bootstrap confidence intervals, collinearity and residual diagnostics, dominance analysis, between- and within-subject analysis of variance, latent class analysis, t-test, z-test, sample size determination), and (7) functions to interact with 'Blimp' and 'Mplus'.

Depends R (>= 4.3.0)

License MIT + file LICENSE

Imports ggplot2, haven, lavaan, lme4, rstudioapi

Suggests boot, clubSandwich, data.table, hdf5r, Matrix, mice, mnormt, mvnmle, nlme, patchwork, readxl, robustlmm, plyr, writexl

Encoding UTF-8

RoxygenNote 7.3.1

NeedsCompilation no

Repository CRAN

Date/Publication 2025-09-06 05:11:28 UTC

Contents

aov.b	4
aov.w	8
blimp	13
blimp.bayes	17
blimp.plot	22
blimp.print	28
blimp.run	31
blimp.update	33
center	37
check.collin	42
check.outlier	45
check.resid	47
chr.color	51
chr.grep	53
chr.gsub	54
chr.omit	56
chr.trim	57
chr.trunc	58
ci.cor	59
ci.mean	73
ci.mean.diff	82
ci.mean.w	85
ci.prop	88
ci.prop.diff	96
ci.var	99
clear	107
cluster.rwg	108
cluster.scores	110
coding	112
coeff.robust	115
coeff.std	120
cohens.d	125
cor.matrix	130
crosstab	134
descript	136
df.check	139
df.duplicated	141
df.head	143
df.long	145
df.merge	148
df.move	150
df.rbind	151
df.rename	153
df.sort	154
df.subset	155
dominance	158

dominance.manual	161
effsize	165
freq	168
indirect	170
item.alpha	174
item.cfa	178
item.invar	187
item.omega	194
item.reverse	198
item.scores	200
lagged	202
libraries	205
mplus	206
mplus.bayes	209
mplus.lca	214
mplus.plot	219
mplus.print	229
mplus.run	235
mplus.update	237
multilevel.cfa	241
multilevel.cor	248
multilevel.descript	253
multilevel.fit	257
multilevel.icc	260
multilevel.indirect	264
multilevel.invar	267
multilevel.omega	273
multilevel.r2	276
multilevel.r2.manual	285
na.as	289
na.auxiliary	291
na.coverage	295
na.descript	297
na.indicator	299
na.pattern	300
na.prop	303
na.satcor	304
na.test	306
plot.misty.object	311
print.misty.object	315
read.data	316
read.dta	318
read.mplus	320
read.sav	321
read.xlsx	323
rec	325
restart	328
result.lca	328

robust.lmer	333
script.copy	335
script.new	337
script.open	338
setsource	340
size.mean	341
skewness	343
summa	346
test.levene	351
test.t	354
test.welch	359
test.z	363
uniq	368
write.data	370
write.dta	372
write.mplus	373
write.result	374
write.sav	376
write.xlsx	379
Index	381

aov.b

Between-Subject Analysis of Variance

Description

This function performs an one-way between-subject analysis of variance (ANOVA) including Tukey HSD post hoc tests for multiple comparison and provides descriptive statistics, effect size measures, and a plot showing bars representing means for each group and error bars for difference-adjusted confidence intervals.

Usage

```
aov.b(formula, data, posthoc = FALSE, conf.level = 0.95, hypo = TRUE,
      descript = TRUE, effsize = FALSE, weighted = FALSE, correct = FALSE,
      digits = 2, p.digits = 3, as.na = NULL, plot = FALSE, bar = TRUE,
      point = FALSE, ci = TRUE, jitter = FALSE, adjust = TRUE,
      point.size = 3, errorbar.width = 0.1, jitter.size = 1.25,
      jitter.width = 0.05, jitter.height = 0, jitter.alpha = 0.1,
      xlab = NULL, ylab = "y", ylim = NULL, ybreaks = ggplot2::waiver(),
      title = NULL, subtitle = "Confidence Interval", filename = NULL,
      width = NA, height = NA, units = c("in", "cm", "mm", "px"), dpi = 600,
      write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

formula	a formula of the form $y \sim \text{group}$ where y is a numeric variable giving the data values and group a numeric variable, character variable or factor with more than two values or factor levels giving the corresponding groups.
data	a matrix or data frame containing the variables in the formula formula.
posthoc	logical: if TRUE, Tukey HSD post hoc test for multiple comparison is conducted.
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
hypo	logical: if TRUE (default), null and alternative hypothesis are shown on the console.
descript	logical: if TRUE (default), descriptive statistics are shown on the console.
effsize	logical: if TRUE, effect size measures η^2 and ω^2 for the ANOVA and Cohen's d for the post hoc tests are shown on the console.
weighted	logical: if TRUE, the weighted pooled standard deviation is used to compute Cohen's d.
correct	logical: if TRUE, correction factor to remove positive bias in small samples is used.
digits	an integer value indicating the number of decimal places to be used for displaying descriptive statistics and confidence interval.
p.digits	an integer value indicating the number of decimal places to be used for displaying the p -value.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
plot	logical: if TRUE, a plot showing error bars for confidence intervals is drawn.
bar	logical: if TRUE (default), bars representing means for each groups are drawn.
point	logical: if TRUE, points representing means for each groups are drawn.
ci	logical: if TRUE (default), error bars representing confidence intervals are drawn.
jitter	logical: if TRUE, jittered data points are drawn.
adjust	logical: if TRUE (default), difference-adjustment for the confidence intervals is applied.
point.size	a numeric value indicating the size aesthetic for the point representing the mean value.
errorbar.width	a numeric value indicating the horizontal bar width of the error bar.
jitter.size	a numeric value indicating the size aesthetic for the jittered data points.
jitter.width	a numeric value indicating the amount of horizontal jitter.
jitter.height	a numeric value indicating the amount of vertical jitter.
jitter.alpha	a numeric value between 0 and 1 for specifying the alpha argument in the <code>geom_histogram</code> function for controlling the opacity of the jittered data points.
xlab	a character string specifying the labels for the x-axis.
ylab	a character string specifying the labels for the y-axis.
ylim	a numeric vector of length two specifying limits of the limits of the y-axis.

ybreaks	a numeric vector specifying the points at which tick-marks are drawn at the y-axis.
title	a character string specifying the text for the title of the plot.
subtitle	a character string specifying the text for the subtitle of the plot.
filename	a character string indicating the filename argument including the file extension in the ggsave function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the filename argument. Note that plots can only be saved when plot = TRUE.
width	a numeric value indicating the width argument (default is the size of the current graphics device) in the ggsave function.
height	a numeric value indicating the height argument (default is the size of the current graphics device) in the ggsave function.
units	a character string indicating the units argument (default is in) in the ggsave function.
dpi	a numeric value indicating the dpi argument (default is 600) in the ggsave function.
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

Post Hoc Test Tukey HSD post hoc test reports Cohen's d based on the non-weighted standard deviation (i.e., weighted = FALSE) when requesting an effect size measure (i.e., effsize = TRUE) following the recommendation by Delacre et al. (2021).

Confidence Intervals Cumming and Finch (2005) pointed out that when 95% confidence intervals (CI) for two separately plotted means overlap, it is still possible that the CI for the difference would not include zero. Baguley (2012) proposed to adjust the width of the CIs by the factor of $\sqrt{2}$ to reflect the correct width of the CI for a mean difference:

$$\hat{\mu}_j \pm t_{n-1, 1-\alpha/2} \frac{\sqrt{2}}{2} \hat{\sigma}_{\hat{\mu}_j}$$

These difference-adjusted CIs around the individual means can be interpreted as if it were a CI for their difference. Note that the width of these intervals is sensitive to differences in the variance and sample size of each sample, i.e., unequal population variances and unequal n alter the interpretation of difference-adjusted CIs.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	data frame with variables used in the current analysis
formula	formula of the current analysis
args	specification of function arguments
plot	ggplot2 object for plotting the results
result	list with result tables, i.e., <code>descript</code> for descriptive statistics, <code>test</code> for the ANOVA table, <code>posthoc</code> for post hoc tests, and <code>aov</code> for the return object of the <code>aov</code> function

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Baguley, T. S. (2012a). *Serious stats: A guide to advanced statistics for the behavioral sciences*. Palgrave Macmillan.
- Cumming, G., and Finch, S. (2005) Inference by eye: Confidence intervals, and how to read pictures of data. *American Psychologist*, 60, 170–80.
- Delacre, M., Lakens, D., Ley, C., Liu, L., & Leys, C. (2021). Why Hedges' g 's based on the non-pooled standard deviation should be reported with Welch's t -test. <https://doi.org/10.31234/osf.io/tu6mp>
- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.

See Also

[aov.w](#), [test.t](#), [test.z](#), [test.levene](#), [test.welch](#), [cohens.d](#), [ci.mean.diff](#), [ci.mean](#)

Examples

```
# Example 1: Between-subject ANOVA
aov.b(mpg ~ gear, data = mtcars)

# Example 2: Between-subject ANOVA
# print effect size measures
aov.b(mpg ~ gear, data = mtcars, effsize = TRUE)

# Example 3: Between-subject ANOVA
# do not print hypotheses and descriptive statistics,
aov.b(mpg ~ gear, data = mtcars, descript = FALSE, hypo = FALSE)

# Example 4: Between-subject ANOVA
# plot results
aov.b(mpg ~ gear, data = mtcars, plot = TRUE)

## Not run:
# Example 5: Write Results into a text file
```

```

aov.b(mpg ~ gear, data = mtcars, write = "ANOVA.txt")

# Example 6: Save plot
aov.b(mpg ~ gear, data = mtcars, plot = TRUE, filename = "Between-Subject_ANOVA.png",
      width = 7, height = 6)

# Example 7: Draw plot in line with the default setting of aov.b()
library(ggplot2)

object <- aov.b(mpg ~ gear, data = mtcars, output = FALSE)

ggplot(object$data, aes(group, y)) +
  geom_bar(stat = "summary", fun = "mean") +
  geom_errorbar(data = ci.mean(object$data, y, group = "group", adjust = TRUE,
                              output = FALSE)$result,
               aes(group, m, ymin = low, ymax = upp), width = 0.1) +
  scale_x_discrete(name = NULL) +
  labs(subtitle = "Two-Sided Difference-Adjusted Confidence Interval") +
  theme_bw() + theme(plot.subtitle = element_text(hjust = 0.5))

## End(Not run)

```

aov.w

Repeated Measures Analysis of Variance (Within-Subject ANOVA)

Description

This function performs an one-way repeated measures analysis of variance (within subject ANOVA) including paired-samples t-tests for multiple comparison and provides descriptive statistics, effect size measures, and a plot showing error bars for difference-adjusted Cousineau-Morey within-subject confidence intervals with jittered data points including subject-specific lines.

Usage

```

aov.w(formula, data, print = c("all", "none", "LB", "GG", "HF"),
      posthoc = FALSE, conf.level = 0.95,
      p.adj = c("none", "bonferroni", "holm", "hochberg", "hommel", "BH", "BY", "fdr"),
      hypo = TRUE, descript = TRUE, epsilon = TRUE, effsize = FALSE, na.omit = TRUE,
      digits = 2, p.digits = 3, as.na = NULL, plot = FALSE, point = TRUE, line = TRUE,
      ci = TRUE, jitter = FALSE, adjust = TRUE, point.size = 3, line.width = 0.5,
      errorbar.width = 0.1, jitter.size = 1.25, jitter.width = 0.05, jitter.alpha = 0.1,
      xlab = NULL, ylab = "y", ylim = NULL, ybreaks = ggplot2::waiver(), title = NULL,
      subtitle = "Confidence Interval", filename = NULL, width = NA, height = NA,
      units = c("in", "cm", "mm", "px"), dpi = 600, write = NULL, append = TRUE,
      check = TRUE, output = TRUE)

```

Arguments

formula	a formula of the form <code>cbind(time1, time2, time3) ~ 1</code> where <code>time1</code> , <code>time2</code> , and <code>time3</code> are numeric variables representing the levels of the within-subject factor, i.e., data are specified in wide-format (i.e., multivariate person level format).
data	a matrix or data frame containing the variables in the formula <code>formula</code> .
print	a character vector indicating which sphericity correction to use, i.e., <code>all</code> for all corrections, <code>none</code> for no correction, <code>LB</code> for lower bound correction, <code>GG</code> for Greenhouse-Geisser correction, and <code>HF</code> , for Huynh-Feldt correction.
posthoc	logical: if <code>TRUE</code> , paired-samples t-tests for multiple comparison are conducted.
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
p.adj	a character string indicating an adjustment method for multiple testing based on p.adjust , i.e., <code>none</code> , <code>bonferroni</code> , <code>holm</code> (default), <code>hochberg</code> , <code>hommel</code> , <code>BH</code> , <code>BY</code> , or <code>fdr</code> .
hypo	logical: if <code>TRUE</code> (default), null and alternative hypothesis are shown on the console.
descript	logical: if <code>TRUE</code> (default), descriptive statistics are shown on the console.
epsilon	logical: if <code>TRUE</code> (default), box indices of sphericity (epsilon) are shown on the console, i.e., lower bound, Greenhouse and Geiser (GG), Huynh and Feldt (HF) and average of GG and HF.
effsize	logical: if <code>TRUE</code> , effect size measures eta-squared (η^2), partial eta-squared (η_p^2), omega-squared (ω^2), and partial omega-squared (ω_p^2) for the repeated measures ANOVA and Cohen's <i>d</i> for the post hoc tests are shown on the console.
na.omit	logical: if <code>TRUE</code> , incomplete cases are removed before conducting the analysis (i.e., listwise deletion).
digits	an integer value indicating the number of decimal places to be used for displaying descriptive statistics and confidence interval.
p.digits	an integer value indicating the number of decimal places to be used for displaying the <i>p</i> -value.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to <code>NA</code> before conducting the analysis.
plot	logical: if <code>TRUE</code> , a plot showing error bars for confidence intervals is drawn.
point	logical: if <code>TRUE</code> (default), points representing means for each groups are drawn.
line	logical: if <code>TRUE</code> (default), a line connecting means of each groups and lines connecting data points are drawn when <code>jitter = TRUE</code> .
ci	logical: if <code>TRUE</code> (default), error bars representing confidence intervals are drawn.
jitter	logical: if <code>TRUE</code> , jittered data points with subject-specific lines are drawn.
adjust	logical: if <code>TRUE</code> (default), difference-adjustment for the Cousineau-Morey within-subject confidence intervals is applied.
point.size	a numeric value indicating the size aesthetic for the point representing the mean value.

<code>line.width</code>	a numeric value indicating the linewidth aesthetic for the line connecting means of each groups.
<code>errorbar.width</code>	a numeric value indicating the horizontal bar width of the error bar.
<code>jitter.size</code>	a numeric value indicating the size aesthetic for the jittered data points.
<code>jitter.width</code>	a numeric value indicating the amount of horizontal jitter.
<code>jitter.alpha</code>	a numeric value between 0 and 1 for specifying the alpha argument in the <code>geom_histogram</code> function for controlling the opacity of the jittered data points.
<code>xlab</code>	a character string specifying the labels for the x-axis.
<code>ylab</code>	a character string specifying the labels for the y-axis.
<code>ylim</code>	a numeric vector of length two specifying limits of the limits of the y-axis.
<code>ybreaks</code>	a numeric vector specifying the points at which tick-marks are drawn at the y-axis.
<code>title</code>	a character string specifying the text for the title for the plot.
<code>subtitle</code>	a character string specifying the text for the subtitle for the plot.
<code>filename</code>	a character string indicating the filename argument including the file extension in the <code>ggsave</code> function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the filename argument. Note that plots can only be saved when <code>plot = TRUE</code> or <code>plot = "boot"</code> .
<code>width</code>	a numeric value indicating the width argument (default is the size of the current graphics device) in the <code>ggsave</code> function.
<code>height</code>	a numeric value indicating the height argument (default is the size of the current graphics device) in the <code>ggsave</code> function.
<code>units</code>	a character string indicating the units argument (default is in) in the <code>ggsave</code> function.
<code>dpi</code>	a numeric value indicating the dpi argument (default is 600) in the <code>ggsave</code> function.
<code>write</code>	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

Sphericity The *F*-Test of the repeated measures ANOVA is based on the assumption of sphericity, which is defined as the assumption that the variance of differences between repeated measures are equal in the population. The Mauchly's test is commonly used to test this hypothesis. However, test of assumptions addresses an irrelevant hypothesis because what matters is the degree of violation rather than its presence (Baguley, 2012a). Moreover, the test is not recommended because it lacks statistical power (Abdi, 2010). Instead, the Box index of sphericity

(ε) should be used to assess the degree of violation of the sphericity assumption. The ε parameter indicates the degree to which the population departs from sphericity with $\varepsilon = 1$ indicating that sphericity holds. As the departure becomes more extreme, ε approaches its lower bound $\hat{\varepsilon}_{lb}$:

$$\hat{\varepsilon}_{lb} = \frac{1}{J - 1}$$

where J is the number of levels of the within-subject factor. Box (1954a, 1954b) suggested a measure for sphericity, which applies to a population covariance matrix. Greenhouse and Geisser (1959) proposed an estimate for ε known as $\hat{\varepsilon}_{gg}$ that can be computed from the sample covariance matrix, whereas Huynh and Feldt (1976) proposed an alternative estimate $\hat{\varepsilon}_{hf}$. These estimates can be used to correct the effect and error df of the F -test. Simulation studies showed that $\hat{\varepsilon}_{gg} \leq \hat{\varepsilon}_{hf}$ and that $\hat{\varepsilon}_{gg}$ tends to be conservative underestimating ε , whereas $\hat{\varepsilon}_{hf}$ tends to be liberal overestimating ε and occasionally exceeding one. Baguley (2012a) recommended to compute the average of the conservative estimate $\hat{\varepsilon}_{gg}$ and the liberal estimate $\hat{\varepsilon}_{hf}$ to assess the sphericity assumption. By default, the function prints results depending on the average $\hat{\varepsilon}_{gg}$ and $\hat{\varepsilon}_{hf}$:

- If the average is less than 0.75 results of the F -Test based on Greenhouse-Geiser correction factor ($\hat{\varepsilon}_{gg}$) is printed.
- If the average is less greater or equal 0.75, but less than 0.95 results of the F -Test based on Huynh-Feldt correction factor ($\hat{\varepsilon}_{hf}$) is printed.
- If the average is greater or equal 0.95 results of the F -Test without any corrections are printed.

Missing Data The function uses listwise deletion by default to deal with missing data. However, the function also allows to use all available observations by conducting the repeated measures ANOVA in long data format when specifying `na.omit = FALSE`. Note that in the presence of missing data, the F -Test without any sphericity corrections may be reliable, but it is not clear whether results based on Greenhouse-Geiser or Huynh-Feldt correction are trustworthy given that pairwise deletion is used for estimating the variance-covariance matrix when computing $\hat{\varepsilon}_{gg}$ and the total number of subjects regardless of missing values (i.e., complete and incomplete cases) are used for computing $\hat{\varepsilon}_{hf}$.

Within-Subject Confidence Intervals The function provides a plot showing error bars for difference-adjusted Cousineau-Morey confidence intervals (Baguley, 2012b). The intervals matches that of a CI for a difference, i.e., non-overlapping CIs corresponds to an inferences of no statistically significant difference. The Cousineau-Morey confidence intervals without adjustment can be used by specifying `adjust = FALSE`.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	list with the data (<code>data</code>) in wide-format (<code>wide</code>), reshaped data in long-format (<code>long</code>), and within-subject confidence intervals (<code>ci</code>)
<code>formula</code>	formula of the current analysis
<code>args</code>	specification of function arguments

<code>plot</code>	ggplot2 object for plotting the results
<code>result</code>	list with result tables, i.e., <code>descript</code> for descriptive statistics, <code>epsilon</code> for a table with indices of sphericity, <code>test</code> for the ANOVA table (none for no sphericity correction, <code>lb</code> for lower bound correction, <code>gg</code> for Greenhouse and Geisser correction, and <code>hf</code> for Huynh and Feldt correction), <code>posthoc</code> for post hoc tests, and <code>aov</code> for the return object of the <code>aov</code> function

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Abdi, H. (2010). The Greenhouse-Geisser correction. In N. J. Salkind (Ed.) *Encyclopedia of Research Design* (pp. 630-634), Sage. <https://dx.doi.org/10.4135/9781412961288>
- Baguley, T. S. (2012a). *Serious stats: A guide to advanced statistics for the behavioral sciences*. Palgrave Macmillan.
- Baguley, T. (2012b). Calculating and graphing within-subject confidence intervals for ANOVA. *Behavior Research Methods*, 44, 158-175. <https://doi.org/10.3758/s13428-011-0123-7>
- Bakerman, R. (2005). Recommended effect size statistics for repeated measures designs. *Behavior Research Methods*, 37, 179-384. <https://doi.org/10.3758/BF03192707>
- Box, G. E. P. (1954a) Some Theorems on Quadratic Forms Applied in the Study of Analysis of Variance Problems, I. Effects of Inequality of Variance in the One-way Classification. *Annals of Mathematical Statistics*, 25, 290–302.
- Box, G. E. P. (1954b) Some Theorems on Quadratic Forms Applied in the Study of Analysis of Variance Problems, II. Effects of Inequality of Variance and of Correlation between Errors in the Two-way Classification. *Annals of Mathematical Statistics*, 25, 484–98.
- Greenhouse, S. W., and Geisser, S. (1959). On methods in the analysis of profile data. *Psychometrika*, 24, 95-112. <https://doi.org/10.1007/BF02289823>
- Huynh, H., and Feldt, L. S. (1976). Estimation of the box correction for degrees of freedom from sample data in randomized block and splitplot designs. *Journal of Educational Statistics*, 1, 69-82. <https://doi.org/10.2307/1164736>
- Olejnik, S., & Algina, J. (2000). Measures of effect size for comparative studies: Applications, interpretations, and limitations. *Contemporary Educational Psychology*, 25, 241-286. <https://doi.org/10.1006/ceps.2000.1040>
- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.

See Also

[aov.b](#), [test.t](#), [test.z](#), [cohens.d](#), [ci.mean.diff](#), [ci.mean](#)

Examples

```
dat <- data.frame(time1 = c(3, 2, 1, 4, 5, 2, 3, 5, 6, 7),
                  time2 = c(4, 3, 6, 5, 8, 6, 7, 3, 4, 5),
                  time3 = c(1, 2, 2, 3, 6, 5, 1, 2, 4, 6))
```

```

# Example 1: Repeated measures ANOVA
aov.w(cbind(time1, time2, time3) ~ 1, data = dat)

# Example 2: Repeated measures ANOVA, print results of all sphericity corrections
aov.w(cbind(time1, time2, time3) ~ 1, data = dat, print = "all")

# Example 3: Repeated measures ANOVA
# print effect size measures
aov.w(cbind(time1, time2, time3) ~ 1, data = dat, effsize = TRUE)

# Example 4: Repeated measures ANOVA, do not print hypotheses and descriptive statistics,
aov.w(cbind(time1, time2, time3) ~ 1, data = dat, descript = FALSE, hypo = FALSE)

# Example 5: Repeated measures ANOVA, plot results
aov.w(cbind(time1, time2, time3) ~ 1, data = dat, plot = TRUE)

## Not run:
# Example 6: Write Results into a text file
aov.w(cbind(time1, time2, time3) ~ 1, data = dat, write = "RM-ANOVA.txt")

# Example 7: Save plot
aov.w(cbind(time1, time2, time3) ~ 1, data = dat, plot = TRUE,
      filename = "Repeated_measures_ANOVA.png", width = 7, height = 6)
## End(Not run)

```

blimp

Create, Run, and Print Blimp Models

Description

This wrapper function creates a Blimp input file, runs the input file by using the `blimp.run()` function, and prints the Blimp output file by using the `blimp.print()` function.

Usage

```

blimp(x, file = "Blimp_Input.imp", data = NULL, comment = FALSE, replace.inp = TRUE,
      blimp.run = TRUE, posterior = FALSE, folder = "Posterior_",
      format = c("csv", "csv2", "excel", "rds", "workspace"), clear = TRUE,
      replace.out = c("always", "never", "modified"), Blimp = .detect.blimp(),
      result = c("all", "default", "algo.options", "data.info", "model.info",
                 "warn.mess", "error.mess", "out.model", "gen.param"),
      exclude = NULL, color = c("none", "blue", "green"),
      style = c("bold", "regular"), not.result = TRUE, write = NULL,
      append = TRUE, check = TRUE, output = TRUE)

```

Arguments

`x` a character string containing the Blimp input text.

file	a character string indicating the name of the Blimp input file with or without the file extension .imp, e.g., "Blimp_Input.imp" or "Blimp_Input.imp".
data	a matrix or data frame from which the variables names for the section VARIABLES are extracted.
comment	logical: if FALSE (default), comments (i.e., text after the # symbol) are removed from the input text specified in the argument x.
replace.inp	logical: if TRUE (default), an existing input file will be replaced.
blimp.run	logical: if TRUE, the input file specified in the argument file containing the input text specified in the argument x is run using the <code>blimp.run()</code> function.
posterior	logical: if TRUE, the posterior distribution including burn-in and post-burn-in phase for all parameters are saved in long format in a file called <code>posterior.*</code> in the folder specified in the argument folder and .imp file name in the format specified in the argument format.
folder	a character string indicating the prefix of the folder for saving the posterior distributions. The default setting is <code>folder = "Posterior_"</code> .
format	a character vector indicating the file format(s) for saving the posterior distributions, i.e., "csv" (default) for <code>write.csv()</code> , "csv2" for <code>write.csv2()</code> , "excel" for <code>write.xlsx()</code> , "rds" for <code>saveRDS()</code> , and "workspace" for <code>write()</code> .
clear	logical: if TRUE (default), the console is cleared after estimating each model.
replace.out	a character string for specifying three settings: "always" (default), which runs all models, regardless of whether an output file for the model exists, "never", which does not run any model that has an existing output file, and "modified", which only runs a model if the modified date for the input file is more recent than the output file modified date.
Blimp	a character string for specifying the name or path of the Blimp executable to be used for running models. This covers situations where Blimp is not in the system's path, or where one wants to test different versions of the Blimp program. Note that there is no need to specify this argument for most users since it has intelligent defaults.
result	a character vector specifying Blimp result sections included in the output (see 'Details' in the <code>blimp.print</code> function).
exclude	a character vector specifying Blimp input command or result sections excluded from the output (see 'Details' in the <code>blimp.print</code> function).
color	a character vector with two elements indicating the colors used for the main headers (e.g., "ALGORITHMIC OPTIONS SPECIFIED: "), and for the headers Outcome Variable: and Missing predictor:, Latent Variable:, and Covariance Matrix:.
style	a character vector with two elements indicating the style used for headers (e.g., "ALGORITHMIC OPTIONS SPECIFIED: "), and for the main headers (e.g., "ALGORITHMIC OPTIONS SPECIFIED: "), and for the headers Outcome Variable: and Missing predictor:, Complete variable:, Latent Variable:, and Covariance Matrix:.
not.result	logical: if TRUE (default), character vector indicating the result sections not requested are shown on the console.
write	a character string naming a file for writing the output into a text file with file extension ".txt" (e.g., "Output.txt").

append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console by using the function <code>blimp.print()</code> .

Details

VARIABLES Section The VARIABLES section used to assign names to the variables in the data set can be specified by using the `data` argument:

- **Write Blimp Data File:** In the first step, the Blimp data file is written by using the `write.mplus()` function, e.g. `write.mplus(data1, file = "data1.dat")`.
- **Specify Blimp Input:** In the second step, the Blimp input is specified as a character string. The VARIABLES option is left out from the Blimp input text, e.g., `input <- 'DATA: data1.dat; \nMODEL: y ~ x1@b1 x2@b2 d2; '`.
- **Run Blimp Input:** In the third step, the Blimp input is run by using the `blimp()` function. The argument `data` needs to be specified given that the VARIABLES section was left out from the Blimp input text in the previous step, e.g., `blimp(input, file = "Ex4.3.imp", data = data1)`.

Note that unlike Mplus, Blimp allows to specify a CSV data file with variable names in the first row. Hence, it is recommended to export the data from R using the `write.csv()` function to specify the data file in the DATA section of the Blimp input file without specifying the VARIABLES section.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>x</code>	a character vector containing the Blimp input text
<code>args</code>	specification of function arguments
<code>write</code>	write command sections
<code>result</code>	list with result sections (<code>result</code>)

Author(s)

Takuya Yanagida

References

Keller, B. T., & Enders, C. K. (2023). *Blimp user's guide* (Version 3). Retrieved from www.appliedmissingdata.com/blimp

See Also

[blimp.update](#), [blimp.run](#), [blimp.print](#), [blimp.plot](#), [blimp.bayes](#)

Examples

```
## Not run:

#-----
# Example 1: Write data, specify input without VARIABLES section, and run input

# Write Data File
# Note that row.names = FALSE needs to be specified
write.csv(data1, file = "data1.csv", row.names = FALSE)

# Specify Blimp input
input1 <- '
DATA: data1.csv;
ORDINAL: d;
MISSING: 999;
FIXED: d;
CENTER: x1 x2;
MODEL: y ~ x1 x2 d;
SEED: 90291;
BURN: 1000;
ITERATIONS: 10000;
'

# Run Blimp input
blimp(input1, file = "Ex4.3.imp")

#-----
# Example 2: Write data, specify input with VARIABLES section, and run input

# Write Data File
write.mplus(data1, file = "data1.dat", input = FALSE)

# Specify Blimp input
input2 <- '
DATA: data1.dat;
VARIABLES: id v1 v2 v3 y x1 d x2 v4;
ORDINAL: d;
MISSING: 999;
FIXED: d;
CENTER: x1 x2;
MODEL: y ~ x1 x2 d;
SEED: 90291;
BURN: 1000;
ITERATIONS: 10000;
'

# Run Blimp input
blimp(input2, file = "Ex4.3.imp")

#-----
# Example 3: Alternative specification using the data argument
```

```
# Write Data File
write.mplus(data1, file = "data1.dat", input = FALSE)

# Specify Blimp input
input3 <- '
DATA: data1.dat;
ORDINAL: d;
MISSING: 999;
FIXED: d;
CENTER: x1 x2;
MODEL: y ~ x1 x2 d;
SEED: 90291;
BURN: 1000;
ITERATIONS: 10000;
'

# Run Blimp input
blimp(input3, file = "Ex4.3.imp", data = data1)

## End(Not run)
```

blimp.bayes

Blimp Summary Measures, Convergence and Efficiency Diagnostics

Description

This function reads the posterior distribution for all parameters saved in long format in a file called `posterior.*` by the function `blimp.run` or `blimp` when specifying `posterior = TRUE` to compute point estimates (i.e., mean, median, and MAP), measures of dispersion (i.e., standard deviation and mean absolute deviation), measures of shape (i.e., skewness and kurtosis), credible intervals (i.e., equal-tailed intervals and highest density interval), convergence and efficiency diagnostics (i.e., potential scale reduction factor R-hat, effective sample size, and Monte Carlo standard error), probability of direction, and probability of being in the region of practical equivalence for the posterior distribution for each parameter. By default, the function computes the maximum of rank-normalized split-R-hat and rank normalized folded-split-R-hat, Bulk effective sample size (Bulk-ESS) for rank-normalized values using split chains, tail effective sample size (Tail-ESS) defined as the minimum of the effective sample size for 0.025 and 0.975 quantiles, the Bulk Monte Carlo standard error (Bulk-MCSE) for the median and Tail Monte Carlo standard error (Tail-MCSE) defined as the maximum of the MCSE for 0.025 and 0.975 quantiles.

Usage

```
blimp.bayes(x, param = NULL,
            print = c("all", "default", "m", "med", "map", "sd", "mad",
                      "skew", "kurt", "eti", "hdi",
                      "rhat", "b.ess", "t.ess", "b.mcse", "t.mcse"),
            m.bulk = FALSE, split = TRUE, rank = TRUE, fold = TRUE,
            pd = FALSE, null = 0, rope = NULL,
            ess.tail = c(0.025, 0.975), mcse.tail = c(0.025, 0.975),
```

```
alternative = c("two.sided", "less", "greater"), conf.level = 0.95,
digits = 2, r.digits = 3, ess.digits = 0, mcse.digits = 3,
p.digits = 3, write = NULL, append = TRUE, check = TRUE,
output = TRUE)
```

Arguments

<code>x</code>	a character string indicating the name of folder containing the posterior.* file, e.g., "Posterior_Ex4.3" or the name of the posterior.* file with or without any file extension, e.g., "Posterior_ExEx4.3/posterior.csv" or "Posterior_ExEx4.3/posterior". Alternatively, a misty object of type blimp can be specified, i.e., result object of the blimp.plot() function. Note that if the posterior file is specified without file extension while multiple posterior.* files in different file formats are available, then the file is read in following order: csv,RData, rds, and xlsx.
<code>param</code>	a numeric vector indicating which parameters to print. Note that the number of the parameter (Param) and the parameter specification (L1, L2, and L3) are provided in the text file "partable.txt".
<code>print</code>	a character vector indicating which summary measures, convergence, and efficiency diagnostics to be printed on the console, i.e. "all" for all summary measures, convergence, and efficiency diagnostics, "m" for the mean, "med" for the median, "MAP" for the maximum a posteriori probability estimate, "med" for the standard deviation, "mad" for the mean absolute deviation, "skew" for the skewness, "kurt" for the kurtosis, "eti" for the equal-tailed credible interval, "hdi" for the highest density credible interval, "rhat" for the potential scale reduction (PSR) factor R-hat convergence diagnostic, "b.ess" for the bulk effective sample size (ESS), "t.ess" for the tail ESS, "b.mcse" for the bulk Monte Carlo standard error (MCSE), and "t.mcse" for the tail MCSE. The default setting is <code>print = c("med", "sd", "skew", "kurt", "eti", "rhat", "b.ess", "t.ess", "b.mcse", "t.mcse")</code> .
<code>m.bulk</code>	logical: if TRUE the Monte Carlo standard error for the mean is computed. The default setting is <code>m.bulk = FALSE</code> , i.e., the Monte Carlo standard error for the median is computed.
<code>split</code>	logical: if TRUE (default), each MCMC chain is split in half before computing R-hat. Note that the argument <code>split</code> is always set to FALSE when computing ESS.
<code>rank</code>	logical: if TRUE (default), rank-normalization is applied to the posterior draws before computing R-hat and ESS. Note that the argument <code>rank</code> is always set to FALSE when computing MCSE.
<code>fold</code>	logical: if TRUE (default), the maximum of rank-normalized split-R-hat and rank normalized folded-split-R-hat is computed. Note that the arguments <code>split</code> and <code>rank</code> are always set to TRUE when specifying <code>fold = TRUE</code> .
<code>pd</code>	logical: if TRUE, the probability of direction is printed on the console.
<code>null</code>	a numeric value considered as a null effect for the probability of direction (default is 0). Note that the value specified in the argument <code>null</code> applies to all parameters which might not be sensible for all parameters.

<code>rope</code>	a numeric vector with two elements indicating the ROPE's lower and upper bounds. ROPE is also depending on the argument <code>alternative</code> , e.g., if <code>rope = c(-0.1, 0.1)</code> , then the actual ROPE is $[-0.1, 0.1]$ given <code>alternative = "two.sided"</code> (default), $[-\infty, 0.1]$ given <code>alternative = "greater"</code> , and $[-0.1, \infty]$ given <code>alternative = "less"</code> . Note that the interval specified in the argument <code>rope</code> applies to all parameters which might not be sensible for all parameters.
<code>ess.tail</code>	a numeric vector with two elements to specify the quantiles for computing the tail ESS. The default setting is <code>tail = c(0.025, 0.975)</code> , i.e., tail ESS is the minimum of effective sample sizes for 0.025 and 0.975 quantiles.
<code>mcse.tail</code>	a numeric vector with two elements to specify the quantiles for computing the tail MCSE. The default setting is <code>tail = c(0.025, 0.975)</code> , i.e., tail MCSE is the maximum of Monte Carlo standard error for 0.025 and 0.975 quantiles.
<code>alternative</code>	a character string specifying the alternative hypothesis for the credible intervals, must be one of <code>"two.sided"</code> (default), <code>"greater"</code> or <code>"less"</code> .
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the credible interval. The default setting is <code>conf.level = 0.95</code> .
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying point estimates, measures of dispersion, and credible intervals.
<code>r.digits</code>	an integer value indicating the number of decimal places to be used for displaying R-hat values.
<code>ess.digits</code>	an integer value indicating the number of decimal places to be used for displaying effective sample sizes.
<code>mcse.digits</code>	an integer value indicating the number of decimal places to be used for displaying Monte Carlo standard errors.
<code>p.digits</code>	an integer value indicating the number of decimal places to be used for displaying the probability of direction and the probability of being in the region of practical equivalence (ROPE).
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension <code>".txt"</code> (e.g., <code>"Output.txt"</code>) or Excel file with file extension <code>".xlsx"</code> (e.g., <code>"Output.xlsx"</code>). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension <code>".txt"</code> specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console by using the function <code>blimp.print()</code> .

Details

Convergence and Efficiency Diagnostics for Markov Chains Convergence and efficiency diagnostics for Markov chains is based on following numeric measures:

- **Potential Scale Reduction (PSR) factor R-hat:** The PSR factor R-hat compares the between- and within-chain variance for a model parameter, i.e., R-hat larger than 1 indicates that the between-chain variance is greater than the within-chain variance and

chains have not mixed well. According to the default setting, the function computes the improved R-hat as recommended by Vehtari et al. (2020) based on rank-normalizing (i.e., `rank = TRUE`) and folding (i.e., `fold = TRUE`) the posterior draws after splitting each MCMC chain in half (i.e., `split = TRUE`). The traditional R-hat used in Blimp can be requested by specifying `split = TRUE`, `rank = FALSE`, and `fold = FALSE`. Note that the traditional R-hat can catch many problems of poor convergence, but fails if the chains have different variances with the same mean parameter or if the chains have infinite variance with one of the chains having a different location parameter to the others (Vehtari et al., 2020). According to Gelman et al. (2014) a R-hat value of 1.1 or smaller for all parameters can be considered evidence for convergence. The Stan Development Team (2024) recommends running at least four chains and a convergence criterion of less than 1.05 for the maximum of rank normalized split-R-hat and rank normalized folded-split-R-hat. Vehtari et al. (2020), however, recommended to only use the posterior samples if R-hat is less than 1.01 because the R-hat can fall below 1.1 well before convergence in some scenarios (Brooks & Gelman, 1998; Vats & Knudon, 2018).

- **Effective Sample Size (ESS):** The ESS is the estimated number of independent samples from the posterior distribution that would lead to the same precision as the autocorrelated samples at hand. According to the default setting, the function computes the ESS based on rank-normalized split-R-hat and within-chain autocorrelation. The function provides the estimated Bulk-ESS (B.ESS) and the Tail-ESS (T.ESS). The Bulk-ESS is a useful measure for sampling efficiency in the bulk of the distribution (i.e., efficiency of the posterior mean), and the Tail-ESS is useful measure for sampling efficiency in the tails of the distribution (e.g., efficiency of tail quantile estimates). Note that by default, the Tail-ESS is the minimum of the effective sample sizes for 2.5% and 97.5% quantiles (`tail = c(0.025, 0.975)`). According to Kruschke (2015), a rank-normalized ESS greater than 400 is usually sufficient to get a stable estimate of the Monte Carlo standard error. However, a ESS of at least 1000 is considered optimal (Zitzmann & Hecht, 2019).
- **Monte Carlo Standard Error (MCSE):** The MCSE is defined as the standard deviation of the chains divided by their effective sample size and reflects uncertainty due to the stochastic algorithm of the Markov Chain Monte Carlo method. The function provides the estimated Bulk-MCSE (B.MCSE) for the margin of error when using the MCMC samples to estimate the posterior mean and the Tail-ESS (T.MCSE) for the margin of error when using the MCMC samples for interval estimation.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>x</code>	a character string indicating the name of the posterior.*
<code>args</code>	specification of function arguments
<code>data</code>	posterior distribution of each parameter estimate in long format
<code>result</code>	result table with summary measures, convergence, and efficiency diagnostics

Note

This function is a modified copy of functions provided in the **rstan** package by Stan Development Team (2024) and **bayestestR** package by Makowski et al. (2019).

Author(s)

Takuya Yanagida

References

- Brooks, S. P. and Gelman, A. (1998). General Methods for Monitoring Convergence of Iterative Simulations. *Journal of Computational and Graphical Statistics*, 7(4): 434–455. MR1665662.
- Gelman, A., & Rubin, D.B. (1992). Inference from iterative simulation using multiple sequences. *Statistical Science*, 7, 457–472. <https://doi.org/10.1214/ss/1177011136>
- Keller, B. T., & Enders, C. K. (2023). *Blimp user's guide* (Version 3). Retrieved from www.appliedmissingdata.com/blimp
- Kruschke, J. (2015). *Doing Bayesian data analysis: A tutorial with R, JAGS, and Stan*. Academic Press.
- Makowski, D., Ben-Shachar, M., & Lüdtke, D. (2019). bayestestR: Describing effects and their uncertainty, existence and significance within the Bayesian framework. *Journal of Open Source Software*, 4(40), 1541. <https://doi.org/10.21105/joss.01541>
- Stan Development Team (2024). *RStan: the R interface to Stan*. R package version 2.32.6. <https://mc-stan.org/>.
- Vats, D. and Knudson, C. (2018). Revisiting the Gelman-Rubin Diagnostic. arXiv:1812.09384.
- Vehtari, A., Gelman, A., Simpson, D., Carpenter, B., & Bürkner, P.-C. (2020). Rank-normalization, folding, and localization: An improved R-hat for assessing convergence of MCMC. *Bayesian analysis*, 16(2), 667–718. <https://doi.org/10.1214/20-BA1221>
- Zitzmann, S., & Hecht, M. (2019). Going beyond convergence in Bayesian estimation: Why precision matters too and how to assess it. *Structural Equation Modeling*, 26(4), 646–661. <https://doi.org/10.1080/10705511.2018.1511111>

See Also

[blimp](#), [blimp.update](#), [blimp.run](#), [blimp.plot](#), [blimp.print](#), [blimp.plot](#),

Examples

```
## Not run:

#-----
# Blimp Example 4.3: Linear Regression

# Example 1a: Default setting, specifying name of the folder
blimp.bayes("Posterior_Ex4.3")

# Example 1b: Default setting, specifying the posterior file
blimp.bayes("Posterior_Ex4.3/posterior.csv")

# Example 2a: Print all summary measures, convergence, and efficiency diagnostics
```

```

blimp.bayes("Posterior_Ex4.3", print = "all")

# Example 3a: Print default measures plus MAP
blimp.bayes("Posterior_Ex4.3", print = c("default", "map"))

# Example 4: Print traditional R-hat in line with Blimp
blimp.bayes("Posterior_Ex4.3", split = TRUE, rank = FALSE, fold = FALSE)

# Example 5: Print probability of direction and the probability of
# being ROPE [-0.1, 0.1]
blimp.bayes("Posterior_Ex4.3", pd = TRUE, rope = c(-0.1, 0.1))

# Example 6: Write Results into a text file
blimp.bayes("Posterior_Ex4.3", write = "Bayes_Summary.txt")

# Example 7b: Write Results into a Excel file
blimp.bayes("Posterior_Ex4.3", write = "Bayes_Summary.xlsx")

## End(Not run)

```

blimp.plot

Blimp Trace Plots and Posterior Distribution Plots

Description

This function reads the posterior distribution including burn-in and post-burn-in phase for all parameters saved in long format in a file called posterior.* by the function blimp.run or blimp when specifying posterior = TRUE to display trace plots and posterior distribution plots.

Usage

```

blimp.plot(x, plot = c("none", "trace", "post"), param = NULL, labels = TRUE,
  burnin = TRUE, point = c("all", "none", "m", "med", "map"),
  ci = c("none", "eti", "hdi"), conf.level = 0.95, hist = TRUE,
  density = TRUE, area = TRUE, alpha = 0.4, fill = "gray85",
  facet.nrow = NULL, facet.ncol = NULL,
  facet.scales = c("fixed", "free", "free_x", "free_y"),
  xlab = NULL, ylab = NULL, xlim = NULL, ylim = NULL,
  xbreaks = ggplot2::waiver(), ybreaks = ggplot2::waiver(),
  xexpand = ggplot2::waiver(), yexpand = ggplot2::waiver(),
  palette = "Set 2", binwidth = NULL, bins = NULL,
  density.col = "#0072B2", shape = 21,
  point.col = c("#CC79A7", "#D55E00", "#009E73"),
  linewidth = 0.6, linetype = "dashed", line.col = "black",
  plot.margin = NULL, legend.title.size = 10, legend.text.size = 10,
  legend.box.margin = NULL, saveplot = c("all", "none", "trace", "post"),
  filename = "Blimp_Plot.pdf", file.plot = c("_TRACE", "_POST"),
  width = NA, height = NA, units = c("in", "cm", "mm", "px"), dpi = 600,
  check = TRUE)

```

Arguments

<code>x</code>	a character string indicating the name of folder containing the <code>posterior.*</code> file, e.g., "Posterior_Ex4.3" or the name of the <code>posterior.*</code> file with or without any file extension, e.g., "Posterior_ExEx4.3/posterior.csv" or "Posterior_ExEx4.3/posterior". Alternatively, a <code>misty</code> object of type <code>blimp</code> can be specified, i.e., result object of the <code>blimp.plot()</code> function. Note that if the <code>posterior</code> file is specified without file extension while multiple <code>posterior.*</code> files in different file formats are available, then the file is read in following order: <code>csv</code> , <code>RData</code> , <code>rds</code> , and <code>xlsx</code> .
<code>plot</code>	a character string indicating the type of plot to display, i.e., "none" for not displaying any plot, "trace" (default) for displaying trace plots, and <code>post</code> for displaying posterior distribution plots.
<code>param</code>	a numeric vector indicating which parameters to print for the trace plots or posterior distribution plots. Note that the number of the parameter (<code>Param</code>) and the parameter specification (<code>L1</code> , <code>L2</code> , and <code>L3</code>) are provided in the text file "partable.txt". Note that parameters with zero variance are excluded by default.
<code>labels</code>	logical: if <code>TRUE</code> (default), parameter labels (e.g., y Beta x for the slope of the regression y on x) are shown in the facet labels. If <code>FALSE</code> , the numbers of the parameter (e.g., Parameter 1 are shown in the the facet labels.
<code>burnin</code>	logical: if <code>FALSE</code> , the burn-in iterations are discarded when displaying trace plots. The default setting for <code>plot = "trace"</code> is <code>TRUE</code> . Note that the burn-in iterations are always discarded when displaying posterior distribution plots (<code>plot = "post"</code>) regardless of the setting of the argument <code>burnin</code> .
<code>point</code>	a character vector indicating the point estimate(s) to be displayed in the posterior distribution plots, i.e., "all" for all point estimates, "none" for not displaying any point estimates, "m" for the posterior mean estimate, "med" (default) for the posterior median estimate, and "map" for the maximum a posterior estimate.
<code>ci</code>	a character string indicating the type of credible interval to be displayed in the posterior distribution plots, i.e., "none" for not displaying any credible intervals, "eti" (default) for displaying the equal-tailed intervals and "hdi" for displaying the highest density interval.
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the credible interval (default is 0.95).
<code>hist</code>	logical: if <code>TRUE</code> (default), histograms are drawn in the posterior probability plots.
<code>density</code>	logical: if <code>TRUE</code> (default), density curves are drawn in the posterior probability plots.
<code>area</code>	logical: if <code>TRUE</code> (default), statistical not significant and statistical significant area is filled with a different color and vertical lines are drawn.
<code>alpha</code>	a numeric value between 0 and 1 for the <code>alpha</code> argument (default is 0.4) for the <code>annotate</code> , and <code>geom_histogram</code> function.
<code>fill</code>	a character string indicating the color for the "fill" argument (default is "gray85") for the <code>annotate</code> and <code>geom_histogram</code> functions.
<code>facet.nrow</code>	a numeric value indicating the <code>nrow</code> argument (default is <code>NULL</code>) for the <code>facet_wrap</code> function.

facet.ncol	a numeric value indicating the ncol argument (default is 2) for the facet_wrap function.
facet.scales	a character string indicating the scales argument (default is "free") for the facet_wrap function.
xlab	a character string indicating the name argument for the scale_x_continuous function.
ylab	a character string indicating the name argument for the scale_y_continuous function.
xlim	a numeric vector with two elements indicating the limits argument (default is NULL) for the scale_x_continuous function.
ylim	a numeric vector with two elements indicating the limits argument (default is NULL) for the scale_y_continuous function.
xbreaks	a numeric vector indicating the breaks argument (default is ggplot2::waiver()) for the scale_x_continuous function.
ybreaks	a numeric vector indicating the breaks argument (default is ggplot2::waiver()) for the scale_y_continuous function.
xexpand	a numeric vector with two elements indicating the expand argument (default is (0.02, 0)) for the scale_x_continuous function.
yexpand	a numeric vector with two elements indicating the expand argument for the scale_y_continuous function. Note that the default setting depends on the type of plot, e.g., (0.02, 0) for the trace plots and expansion(mult = c(0, 0.05)) for the posterior distribution plots.
palette	a character string indicating the palette name (default is "Set 2") for the hcl.colors function. Note that the character string must be one of hcl.pals().
binwidth	a numeric value indicating the binwidth argument (default is to use the number of bins in bins argument) for the geom_histogram function.
bins	a numeric value indicating the bins argument (default is 30) for the geom_histogram function.
density.col	a character string indicating the color argument (default is "#0072B2") for the geom_density function.
shape	a numeric value indicating the shape argument (default is 21) for the geom_point function.
point.col	a character vector with three elements indicating the values argument (default is c("#CC79A7", "#D55E00", "#009E73")) for the scale_color_manual function.
linewidth	a numeric value indicating the linewidth argument (default is 0.6) for the geom_vline function.
linetype	a numeric value indicating the linetype argument (default is "dashed") for the geom_vline function.
line.col	a character string indicating the color argument (default is "black") for the geom_vline function.
plot.margin	a numeric vector indicating the plot.margin argument for the theme function. Note that the default setting depends on the type of the plot, e.g., c(4, 15, -10, 0) for the trace plots, and c(4, 15, 4, 4) for the autocorrelation plots.

legend.title.size	a numeric value indicating the legend.title argument (default is element_text(size = 10)) for the theme function.
legend.text.size	a numeric value indicating the legend.text argument (default is element_text(size = 10)) for the theme function.
legend.box.margin	a numeric vector indicating the legend.box.margin argument for the theme function. Note that the default setting depends on the type of plot, e.g., c(-16, 6, 6, 6) for the trace plots, and c(-25, 6, 6, 6) for the posterior distribution plots with displaying point estimates.
saveplot	a character vector indicating the plot to be saved, i.e., "all" for saving all plots, "none" (default) for not saving any plots, "trace" for saving the trace plots and post for the saving the posterior distribution plots.
filename	a character string indicating the filename argument (default is "Blimp_Plot.pdf") including the file extension for the ggsave function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the filename argument.
file.plot	a character vector with two elements for distinguishing different types of plots. By default, the character string specified in the argument "filename" ("Blimp_Plot") is concatenated with "_TRACE" ("Blimp_Plot_TRACE") for the trace plots, and "_POST" ("Blimp_Plot_POST") for the posterior distribution plots.
width	a numeric value indicating the width argument (default is the size of the current graphics device) for the ggsave function.
height	a numeric value indicating the height argument (default is the size of the current graphics device) for the ggsave function.
units	a character string indicating the units argument (default is in) for the ggsave function.
dpi	a numeric value indicating the dpi argument (default is 600) for the ggsave function.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
x	a character string indicating the name of the posterior.* file
args	specification of function arguments
data	list with posterior distribution of each parameter estimate in long format (plotdat), plot data for the trace plots (trace), and plot data for the posterior distribution plots (post).
plot	list with the trace plots (trace and posterior distribution plots (post))

Author(s)

Takuya Yanagida

ReferencesKeller, B. T., & Enders, C. K. (2023). *Blimp user's guide* (Version 3). Retrieved from www.appliedmissingdata.com/blimp**See Also**[blimp](#), [blimp.update](#), [blimp.run](#), [blimp.print](#), [blimp.plot](#), [blimp.bayes](#)**Examples**

```
## Not run:

#-----
# Blimp Example 4.3: Linear Regression

#.....
# Trace Plots

# Example 1a: Default setting, specifying name of the folder
blimp.plot("Posterior_Ex4.3")

# Example 1b: Default setting, specifying the posterior file
blimp.plot("Posterior_Ex4.3/posterior.csv")

# Example 1c: Print parameters 2, 3, 4, and 5
blimp.plot("Posterior_Ex4.3", param = 2:5)

# Example 1e: Arrange panels in three columns
blimp.plot("Posterior_Ex4.3", ncol = 3)

# Example 1f: Specify "Pastel 1" palette for the hcl.colors function
blimp.plot("Posterior_Ex4.3", palette = "Pastel 1")

#.....
# Posterior Distribution Plots

# Example 2a: Default setting, i.e., posterior median and equal-tailed interval
blimp.plot("Posterior_Ex4.3", plot = "post")

# Example 2b: Display posterior mean and maximum a posteriori
blimp.plot("Posterior_Ex4.3", plot = "post", point = c("m", "map"))

# Example 2c: Display maximum a posteriori and highest density interval
blimp.plot("Posterior_Ex4.3", plot = "post", point = "map", ci = "hdi")

# Example 2d: Do not display any point estimates and credible interval
blimp.plot("Posterior_Ex4.3", plot = "post", point = "none", ci = "none")

# Example 2d: Do not display histograms
```

```

blimp.plot("Posterior_Ex4.3", plot = "post", hist = FALSE)

#.....
# Save Plots

# Example 3a: Save all plots in pdf format
blimp.plot("Posterior_Ex4.3", saveplot = "all")

# Example 3b: Save all plots in png format with 300 dpi
blimp.plot("Posterior_Ex4.3", saveplot = "all", filename = "Blimp_Plot.png", dpi = 300)

# Example 3a: Save posterior distribution plot, specify width and height of the plot
blimp.plot("Posterior_Ex4.3", plot = "none", saveplot = "post",
           width = 7.5, height = 7)

#-----
# Plot from misty.object

# Create misty.object
object <- blimp.plot("Posterior_Ex4.3", plot = "none")

# Trace plot
blimp.plot(object, plot = "trace")

# Posterior distribution plot
blimp.plot(object, plot = "post")

#-----
# Create Plots Manually

# Load ggplot2 package
library(ggplot2)

# Create misty object
object <- blimp.plot("Posterior_Ex4.3", plot = "none")

#.....
# Example 4: Trace Plots

# Extract data
data.trace <- object$data$trace

# Plot
ggplot(data.trace, aes(x = iter, y = value, color = chain)) +
  annotate("rect", xmin = 0, xmax = 1000, ymin = -Inf, ymax = Inf,
         alpha = 0.4, fill = "gray85") +
  geom_line() +
  facet_wrap(~ param, ncol = 2, scales = "free") +
  scale_x_continuous(name = "", expand = c(0.02, 0)) +
  scale_y_continuous(name = "", expand = c(0.02, 0)) +
  scale_colour_manual(name = "Chain",
                    values = hcl.colors(n = 2, palette = "Set 2")) +
  theme_bw() +

```

```

guides(color = guide_legend(nrow = 1, byrow = TRUE)) +
theme(plot.margin = margin(c(4, 15, -10, 0)),
      legend.position = "bottom",
      legend.title = element_text(size = 10),
      legend.text = element_text(size = 10),
      legend.box.margin = margin(c(-16, 6, 6, 6)),
      legend.background = element_rect(fill = "transparent"))

#.....
# Example 5: Posterior Distribution Plots

# Extract data
data.post <- object$data$post

# Plot
ggplot(data.post, aes(x = value)) +
  geom_histogram(aes(y = after_stat(density)), color = "black", alpha = 0.4,
                fill = "gray85") +
  geom_density(color = "#0072B2") +
  geom_vline(data = data.frame(param = levels(data.post$param),
                                stat = tapply(data.post$value, data.post$param, median)),
             aes(xintercept = stat, color = "Median", linewidth = 0.6) +
  geom_vline(data = data.frame(param = levels(data.post$param),
                                low = tapply(data.post$value, data.post$param,
                                              function(y) quantile(y, probs = 0.025))),
             aes(xintercept = low), linetype = "dashed", linewidth = 0.6) +
  geom_vline(data = data.frame(param = levels(data.post$param),
                                upp = tapply(data.post$value, data.post$param,
                                              function(y) quantile(y, probs = 0.975))),
             aes(xintercept = upp), linetype = "dashed", linewidth = 0.6) +
  facet_wrap(~ param, ncol = 2, scales = "free") +
  scale_x_continuous(name = "", expand = c(0.02, 0)) +
  scale_y_continuous(name = "Probability Density, f(x)",
                     expand = expansion(mult = c(0L, 0.05))) +
  scale_color_manual(name = "Point Estimate", values = c(Median = "#D55E00")) +
  labs(caption = "95% Equal-Tailed Interval") +
  theme_bw() +
  theme(plot.margin = margin(c(4, 15, -8, 4)),
        plot.caption = element_text(hjust = 0.5, vjust = 7),
        legend.position = "bottom",
        legend.title = element_text(size = 10),
        legend.text = element_text(size = 10),
        legend.box.margin = margin(c(-30, 6, 6, 6)),
        legend.background = element_rect(fill = "transparent"))

## End(Not run)

```

Description

This function prints the result sections of a Blimp output file (.blimp-out) on the R console. By default, the function prints selected result sections, i.e., Algorithmic Options Specified, Data Information, Model Information, Warning Messages, Outcome Model Estimates, and Generated Parameters.

Usage

```
blimp.print(x,
            result = c("all", "default", "algo.options", "data.info",
                      "model.info", "warn.mess", "error.mess", "out.model", "gen.param"),
            exclude = NULL, color = c("none", "blue", "green"),
            style = c("bold", "regular"), not.result = TRUE,
            write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>x</code>	a character string indicating the name of the Blimp output file with or without the file extension .blimp-out, e.g., "Blimp_Output.blimp-out" or "Blimp_Output". Alternatively, a misty.object of type blimp can be specified, i.e., result object of the blimp.print() function.
<code>result</code>	a character vector specifying Blimp result sections included in the output (see 'Details').
<code>exclude</code>	a character vector specifying Blimp input command or result sections excluded from the output (see 'Details').
<code>color</code>	a character vector with two elements indicating the colors used for the main headers (e.g., "ALGORITHMIC OPTIONS SPECIFIED:"), and for the headers Outcome Variable: and Missing predictor:, Latent Variable:, and Covariance Matrix:.
<code>style</code>	a character vector with two elements indicating the style used for headers (e.g., "ALGORITHMIC OPTIONS SPECIFIED:"), and for the main headers (e.g., "ALGORITHMIC OPTIONS SPECIFIED:"), and for the headers Outcome Variable: and Missing predictor:, Complete variable:, Latent Variable:, and Covariance Matrix:.
<code>not.result</code>	logical: if TRUE (default), character vector indicating the result sections not requested are shown on the console.
<code>write</code>	a character string naming a file for writing the output into a text file with file extension ".txt" (e.g., "Output.txt").
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

Result Sections Following result sections can be selected by using the `result` argument or excluded by using the `exclude` argument:

- "algo.options" for the ALGORITHMIC OPTIONS SPECIFIED section

- "simdat.summary" for the SIMULATED DATA SUMMARIES section
- "order.simdat" for the VARIABLE ORDER IN SIMULATED DATA section
- "burnin.psr" for the BURN-IN POTENTIAL SCALE REDUCTION (PSR) OUTPUT section
- "mh.accept" for the METROPOLIS-HASTINGS ACCEPTANCE RATES section
- "data.info" for the DATA INFORMATION section
- "var.imp" for the VARIABLES IN IMPUTATION MODEL section
- "model.info" for the MODEL INFORMATION section
- "param.label" for the PARAMETER LABELS section
- "warn.mess" for the WARNING MESSAGES section
- "fit" for the MODEL FIT section
- "cor.resid" for the CORRELATIONS AMONG RESIDUALS section
- "out.model" for the OUTCOME MODEL ESTIMATES section
- "pred.model" for the PREDICTOR MODEL ESTIMATES section
- "gen.param" for the GENERATED PARAMETERS section
- "order.impdat" for the VARIABLE ORDER IN IMPUTED DATA section

Note that all result sections are requested by specifying `result = "all"`. The `result` argument is also used to select one (e.g., `result = "algo.options"`) or more than one result sections (e.g., `result = c("algo.options", "fit")`), or to request result sections in addition to the default setting (e.g., `result = c("default", "fit")`). The `exclude` argument is used to exclude result sections from the output (e.g., `exclude = "algo.options"`).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>x</code>	character string or <code>misty</code> object
<code>args</code>	specification of function arguments
<code>print</code>	print objects
<code>notprint</code>	character vectors indicating the result sections not requested
<code>result</code>	list with Blimp version (<code>blimp</code>) and result sections (<code>result</code>)

Author(s)

Takuya Yanagida

References

Keller, B. T., & Enders, C. K. (2023). *Blimp user's guide* (Version 3). Retrieved from www.appliedmissingdata.com/blimp

See Also

[blimp](#), [blimp.update](#), [blimp.run](#), [blimp.plot](#), [blimp.bayes](#)

Examples

```
## Not run:

#-----
# Blimp Example 4.3: Linear Regression

# Example 1a: Default setting
blimp.print("Ex4.3.blimp-out")

# Example 1c: Print OUTCOME MODEL ESTIMATES only
blimp.print("Ex4.3.blimp-out", result = "out.model")

# Example 1d: Print MODEL FIT in addition to the default setting
blimp.print("Ex4.3.blimp-out", result = c("default", "fit"))

# Example 1e: Exclude DATA INFORMATION section
blimp.print("Ex4.3.blimp-out", exclude = "data.info")

# Example 1f: Print all result sections, but exclude MODEL FIT section
blimp.print("Ex4.3.blimp-out", result = "all", exclude = "fit")

# Example 1g: Print result section in a different order
blimp.print("Ex4.3.blimp-out", result = c("model.info", "fit", "algo.options"))

#-----
# misty.object of type 'blimp.print'

# Example 2
# Create misty.object
object <- blimp.print("Ex4.3.blimp-out", output = FALSE)

# Print misty.object
blimp.print(object)

#-----
# Write Results

# Example 3: Write Results into a text file
blimp.print("Ex4.3.blimp-out", write = "Output_4-3.txt")

## End(Not run)
```

blimp.run

Run Blimp Models

Description

This function runs a group of Blimp models (.imp files) located within a single directory or nested within subdirectories.

Usage

```
blimp.run(target = getwd(), recursive = FALSE,
          replace.out = c("always", "never", "modified"), posterior = FALSE,
          folder = "Posterior_", format = c("csv", "csv2", "xlsx", "rds", "RData"),
          clear = FALSE, Blimp = .detect.blimp(), check = TRUE)
```

Arguments

target	a character string indicating the directory containing Blimp input files (.imp) to run, a character string indicating a single .imp file to run, or a character vector for multiple .imp files to run. May be a full path, relative path, a file name, or a vector of file names within the working directory.
recursive	logical: if TRUE, run all models nested in subdirectories within a directory. Not relevant if a single or multiple .imp files were specified for the argument target.
replace.out	a character string for specifying three settings: "always" (default), which runs all models, regardless of whether an output file for the model exists, "never", which does not run any model that has an existing output file, and "modified", which only runs a model if the modified date for the input file is more recent than the output file modified date.
posterior	logical: if TRUE, the posterior distribution including burn-in and post-burn-in phase for all parameters are saved in long format in a file called posterior.* in the folder specified in the argument folder and .imp file name in the format specified in the argument format.
folder	a character string indicating the prefix of the folder for saving the posterior distributions. The default setting is folder = "Posterior_".
format	a character vector indicating the file format(s) for saving the posterior distributions, i.e., "csv" (default) for write.csv(), "csv2" for write.csv2(), "xlsx" for write.xlsx(), "rds" for saveRDS(), and "RData" for write().
clear	logical: if TRUE, the console is cleared after estimating each model.
Blimp	a character string for specifying the name or path of the Blimp executable to be used for running models. This covers situations where Blimp is not in the system's path, or where one wants to test different versions of the Blimp program. Note that there is no need to specify this argument for most users since it has intelligent defaults.
check	logical: if TRUE (default), argument specification is checked.

Value

None.

Note

This function is based on the detect_blimp() and rblimp() function in the **rblimp** package by Brian T.Keller (2024).

Author(s)

Takuya Yanagida

References

Keller, B. T., & Enders, C. K. (2023). *Blimp user's guide* (Version 3). Retrieved from www.appliedmissingdata.com/blimp

Keller B (2024). *rblimp: Integration of Blimp Software into R*. R package version 0.1.31. <https://github.com/blimp-stats/rblimp>

See Also

[blimp](#), [blimp.update](#), [blimp.print](#), [blimp.plot](#), [blimp.bayes](#)

Examples

```
## Not run:

# Example 1: Run Blimp models located within the current working directory
blimp.run()

# Example 2: Run Blimp models located nested within subdirectories
blimp.run(recursive = TRUE)

# Example 3: Run Blimp input file
blimp.run("Ex4.1a.imp")

# Example 4: Run Blimp input files
blimp.run(c("Ex4.1a.imp", "Ex4.1b.imp"))

# Example 5: Run Blimp models, save posterior distribution in a R workspace
blimp.run(posterior = TRUE, format = "workspace")

## End(Not run)
```

blimp.update

Blimp Input Updating

Description

This function updates specific input command sections of a `misty` object of type `blimp` to create an updated Blimp input file, run the updated input file by using the `blimp.run()` function, and print the updated Blimp output file by using the `blimp.print()` function.

Usage

```
blimp.update(x, update, file = "Blimp_Input_Update.imp", comment = FALSE,
  replace.inp = TRUE, blimp.run = TRUE, posterior = FALSE,
  folder = "Posterior_",
  format = c("csv", "csv2", "xlsx", "rds", "RData"),
  clear = TRUE, replace.out = c("always", "never", "modified"),
  Blimp = .detect.blimp(),
  result = c("all", "default", "algo.options", "data.info",
    "model.info", "warn.mess", "out.model", "gen.param"),
  exclude = NULL, color = c("none", "blue", "violet"),
  style = c("bold", "regular"), not.result = TRUE,
  write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>x</code>	misty.object object of type blimp.
<code>update</code>	a character vector containing the updated input command sections.
<code>file</code>	a character string indicating the name of the updated Blimp input file with or without the file extension .imp, e.g., "Blimp_Input_Update.imp" or "Blimp_Input_Update.imp".
<code>comment</code>	logical: if FALSE (default), comments (i.e., text after the # symbol) are removed from the input text specified in the argument <code>x</code> .
<code>replace.inp</code>	logical: if TRUE (default), an existing input file will be replaced.
<code>blimp.run</code>	logical: if TRUE, the input file specified in the argument <code>file</code> containing the input text specified in the argument <code>x</code> is run using the <code>blimp.run()</code> function.
<code>posterior</code>	logical: if TRUE, the posterior distribution including burn-in and post-burn-in phase for all parameters are saved in long format in a file called <code>posterior.*</code> in the folder specified in the argument <code>folder</code> and .imp file name in the format specified in the argument <code>format</code> .
<code>folder</code>	a character string indicating the prefix of the folder for saving the posterior distributions. The default setting is <code>folder = "Posterior_"</code> .
<code>format</code>	a character vector indicating the file format(s) for saving the posterior distributions, i.e., "csv" (default) for <code>write.csv()</code> , "csv2" for <code>write.csv2()</code> , "xlsx" for <code>write.xlsx()</code> , "rds" for <code>saveRDS()</code> , and "RData" for <code>write()</code> .
<code>clear</code>	logical: if TRUE (default), the console is cleared after estimating each model.
<code>replace.out</code>	a character string for specifying three settings: "always" (default), which runs all models, regardless of whether an output file for the model exists, "never", which does not run any model that has an existing output file, and "modified", which only runs a model if the modified date for the input file is more recent than the output file modified date.
<code>Blimp</code>	a character string for specifying the name or path of the Blimp executable to be used for running models. This covers situations where Blimp is not in the system's path, or where one wants to test different versions of the Blimp program. Note that there is no need to specify this argument for most users since it has intelligent defaults.

result	a character vector specifying Blimp result sections included in the output (see 'Details' in the <code>blimp.print</code> function).
exclude	a character vector specifying Blimp input command or result sections excluded from the output (see 'Details' in the <code>blimp.print</code> function).
color	a character vector with two elements indicating the colors used for headers (e.g., "ALGORITHMIC OPTIONS SPECIFIED: "), and for the header Outcome Variable: and Missing predictor: including variables names.
style	a character vector with two elements indicating the style used for headers (e.g., "ALGORITHMIC OPTIONS SPECIFIED: "), and for the header Outcome Variable: and Missing predictor: including variables names, i.e., regular, for regular text, bold for bold text, italic, for italic text, and underline for underline text.
not.result	logical: if TRUE (default), character vector indicating the result sections not requested are shown on the console.
write	a character string naming a file for writing the output into a text file with file extension ".txt" (e.g., "Output.txt").
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console by using the function <code>blimp.print()</code> .
data	a matrix or data frame from which the variables names for the section VARIABLES are extracted when using the ... specification in the VARIABLES section.

Details

Bimp Input Sections The function is used to update following Blimp input sections:

- DATA
- VARIABLES
- CLUSTERID
- ORDINAL
- NOMINAL
- COUNT
- WEIGHT
- MISSING
- LATENT
- RANDEFFECT
- TRANSFORM
- BYGROUP
- FIXED
- CENTER
- MODEL
- SIMPLE

- PARAMETERS
- TEST
- FCS
- SIMUALTE
- SEED
- BURN
- ITERATIONS
- CHAINS
- NIMPS
- THIN
- OPTIONS
- OUTPUT
- SAVE

The –; Specification The ---; specification is used to remove entire sections (e.g., CENTER: ---;) from the Blimp input. Note that ---; including the semicolon ; needs to be specified, i.e., --- without the semicolon ; will result in an error message.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>x</code>	<code>misty.object</code> object of type <code>blimp</code>
<code>update</code>	a character vector containing the updated Blimp input command sections
<code>args</code>	specification of function arguments
<code>write</code>	updated write command sections
<code>result</code>	list with result sections (<code>result</code>)

Author(s)

Takuya Yanagida

References

Keller, B. T., & Enders, C. K. (2023). *Blimp user's guide* (Version 3). Retrieved from www.appliedmissingdata.com/blimp

See Also

[blimp.run](#), [blimp.print](#), [blimp.plot](#), [blimp.bayes](#)

Examples

```
## Not run:

#-----
# Example 1a: Update BURN and ITERATIONS section

# Specify Blimp input
input <- '
DATA: data1.csv;
ORDINAL: d;
MISSING: 999;
FIXED: d;
CENTER: x1 x2;
MODEL: y ~ x1 x2 d;
SEED: 90291;
BURN: 1000;
ITERATIONS: 10000;
'

# Run Blimp input
mod0 <- blimp(input, file = "Ex4.3.imp", clear = FALSE)

# Update sections
update1 <- '
BURN: 5000;
ITERATIONS: 20000;
'

# Run updated Blimp input
mod1 <- blimp.update(mod0, update1, file = "Ex4.3_update1.imp")

#-----
# Example 1b: Remove CENTER section

# Remove section
update2 <- '
CENTER: ---;
'

# Run updated Blimp input
mod2 <- blimp.update(mod1, update2, file = "Ex4.3_update2.imp")

## End(Not run)
```

center

Centering Predictor Variables in Single-Level and Multilevel Data

Description

This function centers predictor variables in single-level data, two-level data, and three-level data at the grand mean (CGM, i.e., grand mean centering) or within cluster (CWC, i.e., group mean

centering).

Usage

```
center(data, ..., cluster = NULL, type = c("CGM", "CWC"),
       cwc.mean = c("L2", "L3"), value = NULL, append = TRUE, name = ".c",
       as.na = NULL, check = TRUE)
```

Arguments

<code>data</code>	a numeric vector for centering a predictor variable, or a data frame for centering more than one predictor variable.
<code>...</code>	an expression indicating the variable names in data e.g., <code>center(dat, x1, x2)</code> for centering the variables <code>x1</code> and <code>x2</code> in the data frame <code>dat</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>cluster</code>	a character string indicating the name of the cluster variable in data for a two-level model, a character vector indicating the names of the cluster variables in data for a three-level model, or a vector or data frame representing the nested grouping structure (i.e., group or cluster variables). Alternatively, a character string or character vector indicating the variable name(s) of the cluster variable(s) in data. Note that the cluster variable at Level 3 come first in a three-level model, i.e., <code>cluster = c("level3", "level2")</code> .
<code>type</code>	a character string indicating the type of centering, i.e., "CGM" for centering at the grand mean (i.e., grand mean centering, default when <code>cluster = NULL</code>) or "CWC" for centering within cluster (i.e., group mean centering, default when specifying the argument <code>cluster</code>).
<code>cwc.mean</code>	a character string indicating the type of centering of a level-1 predictor variable in a three-level model, i.e., L2 (default) for centering the predictor variable at the level-2 cluster means, and L3 for centering the predictor variable at the level-3 cluster means.
<code>value</code>	a numeric value for centering on a specific user-defined value. Note that this option is only available when specifying a single-level predictor variable, i.e., <code>cluster = NULL</code> .
<code>append</code>	logical: if TRUE (default), centered variable(s) are appended to the data frame specified in the argument <code>data</code> .
<code>name</code>	a character string or character vector indicating the names of the centered predictor variables. By default, centered predictor variables are named with the ending ".c" resulting in e.g. "x1.c" and "x2.c". Variable names can also be specified by using a character vector matching the number of variables (e.g., <code>name = c("center.x1", "center.x2")</code>).
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to data but not to cluster.
<code>check</code>	logical: if TRUE (default), argument specification is checked.

Details

Single-Level Data Predictor variables in single-level data can only be centered at the grand mean (CGM) by specifying type = "CGM":

$$x_i - \bar{x}_.$$

where x_i is the predictor value of observation i and $\bar{x}_.$ is the average x score. Note that predictor variables can be centered on any meaningful value specifying the argument value, e.g., a predictor variable centered at 5 by applying following formula:

$$x_i - \bar{x}_. + 5$$

resulting in a mean of the centered predictor variable of 5.

Two-Level Data Level-1 (L1) predictor variables in two-level data can be centered at the grand mean (CGM) by specifying type = "CGM":

$$x_{ij} - \bar{x}_{..}$$

where x_{ij} is the predictor value of observation i in L2 cluster j and $\bar{x}_{..}$ is the average x score. L1 predictor variables are centered at the group mean (CWC) by specifying type = "CWC" (Default):

$$x_{ij} - \bar{x}_{.j}$$

where $\bar{x}_{.j}$ is the average x score in cluster j .

Level-2 (L1) predictor variables in two-level data can only be centered at the grand mean:

$$x_{.j} - \bar{x}_{..}$$

where $x_{.j}$ is the predictor value of Level 2 cluster j and $\bar{x}_{..}$ is the average Level-2 cluster score. Note that the cluster membership variable needs to be specified when centering a L2 predictor variable in two-level data. Otherwise the average x_{ij} individual score instead of the average $x_{.j}$ cluster score is used to center the predictor variable.

Three-Level Data Level-1 (L1) predictor variables in three-level data can be centered at the grand mean (CGM) by specifying type = "CGM":

$$x_{ijk} - \bar{x}_{...}$$

where x_{ijk} is the predictor value of observation i in Level-2 cluster j within Level-3 cluster k and $\bar{x}_{...}$ is the average x score.

L1 predictor variables are centered within cluster (CWC) by specifying type = "CWC" (Default). However, L1 predictor variables can be either centered within Level-2 clusters (cwc.mean = "L2", Default, see Brincks et al., 2017):

$$x_{ijk} - \bar{x}_{.jk}$$

or within Level-3 clusters (cwc.mean = "L3", see Enders, 2013):

$$x_{ijk} - \bar{x}_{..k}$$

where $\bar{x}_{.jk}$ is the average x score in Level-2 cluster j within Level-3 cluster k and $\bar{x}_{..k}$ is the average x score in Level-3 cluster k .

Level-2 (L2) predictor variables in three-level data can be centered at the grand mean (CGM) by specifying type = "CGM":

$$x_{.jk} - \bar{x}_{...}$$

where $x_{.jk}$ is the predictor value of Level-2 cluster j within Level-3 cluster k and $\bar{x}_{...}$ is the average Level-2 cluster score.

L2 predictor variables are centered within cluster (CWC) by specifying type = "CWC" (Default):

$$x_{.jk} - \bar{x}_{..k}$$

where $\bar{x}_{..k}$ is the average x score in Level-3 cluster k .

Level-3 (L3) predictor variables in three-level data can only be centered at the grand mean:

$$x_{..k} - \bar{x}_{...}$$

where $x_{..k}$ is the predictor value of Level-3 cluster k and $\bar{x}_{...}$ is the average Level-3 cluster score. Note that the cluster membership variables at Level 2 and Level 3 need to be specified when centering a L3 predictor variable in three-level data.

Value

Returns a numeric vector or data frame with the same length or same number of rows as data containing the centered variable(s).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Brincks, A. M., Enders, C. K., Llabre, M. M., Bulotsky-Shearer, R. J., Prado, G., & Feaster, D. J. (2017). Centering predictor variables in three-level contextual models. *Multivariate Behavioral Research*, 52(2), 149–163. <https://doi.org/10.1080/00273171.2016.1256753>
- Chang, C.-N., & Kwok, O.-M. (2022) Partitioning Variance for a Within-Level Predictor in Multi-level Models. *Structural Equation Modeling: A Multidisciplinary Journal*. Advance online publication. <https://doi.org/10.1080/10705511.2022.2051175>
- Enders, C. K. (2013). Centering predictors and contextual effects. In M. A. Scott, J. S. Simonoff, & B. D. Marx (Eds.), *The Sage handbook of multilevel modeling* (pp. 89-109). Sage. <https://dx.doi.org/10.4135/9781446247600>
- Enders, C. K., & Tofighi, D. (2007). Centering predictor variables in cross-sectional multilevel models: A new look at an old issue. *Psychological Methods*, 12, 121-138. <https://doi.org/10.1037/1082-989X.12.2.121>
- Rights, J. D., Preacher, K. J., & Cole, D. A. (2020). The danger of conflating level-specific effects of control variables when primary interest lies in level-2 effects. *British Journal of Mathematical & Statistical Psychology*, 73, 194-211. <https://doi.org/10.1111/bmsp.12194>

Yaremych, H. E., Preacher, K. J., & Hedeker, D. (2021). Centering categorical predictors in multi-level models: Best practices and interpretation. *Psychological Methods*. Advance online publication. <https://doi.org/10.1037/met0000434>

See Also

[coding](#), [cluster.scores](#), [rec](#), [item.reverse](#), [cluster.rwg](#), [item.scores](#).

Examples

```
#-----
# Predictor Variables in Single-Level Data

# Example 1a: Center predictor 'disp' at the grand mean
center(mtcars, disp, append = FALSE)

# Alternative specification without using the '...' argument
center(mtcars$disp)

# Example 1b: Center predictors 'disp' and 'hp' at the grand mean and append to 'mtcars'
center(mtcars, disp, hp)

# Alternative specification without using the '...' argument
cbind(mtcars, center(mtcars[, c("disp", "hp")]))

# Example 1c: Center predictor 'disp' at the value 3
center(mtcars, disp, value = 3)

# Example 1d: Center predictors 'disp' and 'hp' and label with the suffix ".v"
center(mtcars, disp, hp, name = ".v")

#-----
# Predictor Variables in Two-Level Data

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

# Example 2a: Center L1 predictor 'y1' within cluster
center(Demo.twolevel, y1, cluster = "cluster")

# Alternative specification without using the '...' argument
center(Demo.twolevel$y1, cluster = Demo.twolevel$cluster)

# Example 2b: Center L2 predictor 'w2' at the grand mean
center(Demo.twolevel, w1, cluster = "cluster")

# Example 2c: Center L1 predictor 'y1' within cluster and L2 predictor 'w1' at the grand mean
center(Demo.twolevel, y1, w1, cluster = "cluster")

#-----
# Predictor Variables in Three-Level Data
```

```
# Create arbitrary three-level data
Demo.threelevel <- data.frame(Demo.twolevel, cluster2 = Demo.twolevel$cluster,
                             cluster3 = rep(1:10, each = 250))

# Example 3a: Center L1 predictor 'y1' within L2 cluster
center(Demo.threelevel, y1, cluster = c("cluster3", "cluster2"))

# Example 3b: Center L1 predictor 'y1' within L3 cluster
center(Demo.threelevel, y1, cluster = c("cluster3", "cluster2"), cwc.mean = "L3")

# Example 3c: Center L1 predictor 'y1' within L2 cluster and L2 predictor 'w1' within L3 cluster
center(Demo.threelevel, y1, w1, cluster = c("cluster3", "cluster2"))
```

check.collin	<i>Collinearity Diagnostics</i>
--------------	---------------------------------

Description

This function computes tolerance, standard error inflation factor, variance inflation factor, eigenvalues, condition index, and variance proportions for linear, generalized linear, and mixed-effects models.

Usage

```
check.collin(model, print = c("all", "vif", "eigen"), digits = 3, p.digits = 3,
             write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

model	a fitted model of class "lm", "glm", "lmerMod", "lmerModLmerTest", "glmerMod", "lme", or "glmmTMB".
print	a character vector indicating which results to show, i.e. "all", for all results, "vif" for tolerance, std. error inflation factor, and variance inflation factor, or eigen for eigenvalue, condition index, and variance proportions.
digits	an integer value indicating the number of decimal places to be used for displaying results.
p.digits	an integer value indicating the number of decimal places to be used for displaying the <i>p</i> -value.
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

Collinearity diagnostics can be conducted for objects returned from the `lm()` and `glm()` function, but also from objects returned from the `lmer()` and `glmer()` function from the **lme4** package, `lme()` function from the **nlme** package, and the `glmmTMB()` function from the **glmmTMB** package.

The generalized variance inflation factor (Fox & Monette, 1992) is computed for terms with more than 1 df resulting from factors with more than two levels. The generalized VIF (GVIF) is interpretable as the inflation in size of the confidence ellipse or ellipsoid for the coefficients of the term in comparison with what would be obtained for orthogonal data. GVIF is invariant to the coding of the terms in the model. In order to adjust for the dimension of the confidence ellipsoid, $GVIF^{\frac{1}{2df}}$ is computed. Note that the adjusted GVIF (aGVIF) is actually a generalized standard error inflation factor (GSIF). Thus, the aGVIF needs to be squared before applying a common cutoff threshold for the VIF (e.g., $VIF > 10$). Note that the output of `check.collin()` function reports either the variance inflation factor or the squared generalized variance inflation factor in the column VIF, while the standard error inflation factor or the adjusted generalized variance inflation factor is reported in the column SIF.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>model</code>	model specified in the <code>model</code> argument
<code>args</code>	specification of function arguments
<code>result</code>	list with result tables, i.e., <code>coef</code> for the regression table including tolerance, std. error inflation factor and variance inflation factors, <code>vif</code> for the tolerance, std. error inflation factor, and variance inflation factor, and <code>eigen</code> for eigenvalue condition index, and variance proportion

Note

The computation of the VIF and the GVIF is based on the `vif()` function in the **car** package by John Fox, Sanford Weisberg and Brad Price (2020), and the computation of eigenvalues, condition index, and variance proportions is based on the `ols_eigen_cindex()` function in the **olsrr** package by Aravind Hebbali (2020).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Fox, J., & Monette, G. (1992). Generalized collinearity diagnostics. *Journal of the American Statistical Association*, 87, 178-183.
- Fox, J., Weisberg, S., & Price, B. (2020). *car: Companion to Applied Regression*. R package version 3.0-8. <https://cran.r-project.org/web/packages/car/>
- Hebbali, A. (2020). *olsrr: Tools for building OLS regression models*. R package version 0.5.3. <https://cran.r-project.org/web/packages/olsrr/>

See Also

`check.outlier, lm`

Examples

```
dat <- data.frame(group = c(1, 1, 1, 2, 2, 2, 3, 3, 3, 4, 4, 4),
  x1 = c(3, 2, 4, 9, 5, 3, 6, 4, 5, 6, 3, 5),
  x2 = c(1, 4, 3, 1, 2, 4, 3, 5, 1, 7, 8, 7),
  x3 = c(7, 3, 4, 2, 5, 6, 4, 2, 3, 5, 2, 8),
  x4 = c("a", "b", "a", "c", "c", "c", "a", "b", "b", "c", "a", "c"),
  y1 = c(2, 7, 4, 4, 7, 8, 4, 2, 5, 1, 3, 8),
  y2 = c(0, 1, 0, 1, 1, 1, 0, 0, 0, 1, 0, 1), stringsAsFactors = TRUE)

#-----
# Linear model

# Estimate linear model with continuous predictors
mod.lm1 <- lm(y1 ~ x1 + x2 + x3, data = dat)

# Example 1: Tolerance, std. error, and variance inflation factor
check.collin(mod.lm1)

# Example 2: Tolerance, std. error, and variance inflation factor
# Eigenvalue, Condition index, and variance proportions
check.collin(mod.lm1, print = "all")

# Estimate model with continuous and categorical predictors
mod.lm2 <- lm(y1 ~ x1 + x2 + x3 + x4, data = dat)

# Example 3: Tolerance, generalized std. error, and variance inflation factor
check.collin(mod.lm2)

#-----
# Generalized linear model

# Estimate logistic regression model with continuous predictors
mod.glm <- glm(y2 ~ x1 + x2 + x3, data = dat, family = "binomial")

# Example 4: Tolerance, std. error, and variance inflation factor
check.collin(mod.glm)

## Not run:
#-----
# Linear mixed-effects model

# Load lme4, nlme, and glmmTMB package
libraries(lme4, nlme, glmmTMB)

# Estimate linear mixed-effects model using lme4 package
mod.lmer <- lmer(y1 ~ x1 + x2 + x3 + (1|group), data = dat)

# Example 5: Tolerance, std. error, and variance inflation factor
```

```

check.collin(mod.lmer)

# Estimate linear mixed-effects model using nlme package
mod.lme <- lme(y1 ~ x1 + x2 + x3, random = ~ 1 | group, data = dat)

# Example 6: Tolerance, std. error, and variance inflation factor
check.collin(mod.lme)

# Estimate linear mixed-effects model using glmmTMB package
mod.glmmTMB1 <- glmmTMB(y1 ~ x1 + x2 + x3 + (1|group), data = dat)

# Example 7: Tolerance, std. error, and variance inflation factor
check.collin(mod.glmmTMB1)

#-----
# Generalized linear mixed-effects model

# Estimate mixed-effects logistic regression model using lme4 package
mod.glmer <- glmer(y2 ~ x1 + x2 + x3 + (1|group), data = dat, family = "binomial")

# Example 8: Tolerance, std. error, and variance inflation factor
check.collin(mod.glmer)

# Estimate mixed-effects logistic regression model using glmmTMB package
mod.glmmTMB2 <- glmmTMB(y2 ~ x1 + x2 + x3 + (1|group), data = dat, family = "binomial")

# Example 9: Tolerance, std. error, and variance inflation factor
check.collin(mod.glmmTMB2)

#-----
# Write Results

# Example 10: Write Results into a text file
check.collin(mod.lm1, write = "Diagnostics.txt")

## End(Not run)

```

check.outlier

Statistical Measures for Leverage, Distance, and Influence

Description

This function computes statistical measures for leverage, distance, and influence for linear models estimated by using the `lm()` function. Mahalanobis distance and hat values are computed for quantifying *leverage*, standardized leverage-corrected residuals and studentized leverage-corrected residuals are computed for quantifying *distance*, and Cook's distance and DfBetas are computed for quantifying *influence*.

Usage

```
check.outlier(model, append = TRUE, check = TRUE, ...)
```

Arguments

model	a fitted model of class "lm".
append	logical: if TRUE (default), statistical measures for leverage, distance, and influence are appended to the data frame in model\$model.
check	logical: if TRUE (default), argument specification is checked.
...	further arguments to be passed to or from methods.

Details

In regression analysis, an observation can be extreme in three major ways (see Darlington & Hayes, p. 484): (1) An observation has high **leverage** if it has a atypical pattern of values on the predictors, (2) an observation has high **distance** if its observed outcome value Y_i has a large deviation from the predicted value $\hat{Y}_i$, and (3) an observation has high **influence** if its inclusion substantially changes the estimates for the intercept and/or slopes.

Value

Returns a data frame with following entries:

idout	ID variable
mahal	Mahalanobis distance
hat	hat values
rstand	standardized leverage-corrected residuals
rstud	studentized leverage-corrected residuals
cook	Cook's distance
Intercept.dfb	DFBetas for the intercept
pred1.dfb	DFBetas for the slope of the predictor pred1
....dfb	DFBetas for the slope of the predictor ...

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Darlington, R. B., & Hayes, A. F. (2017). *Regression analysis and linear models: Concepts, applications, and implementation*. The Guilford Press.

See Also

[check.collin, lm](#)

Examples

```
# Example 1: Statistical measures for leverage, distance, and influence
check.outlier(lm(mpg ~ cyl + disp + hp, data = mtcars))

# Example 2: Append statistical measures to the mtcars data frame
cbind(mtcars,
      check.outlier(lm(mpg ~ cyl + disp + hp, data = mtcars), append = FALSE))
```

check.resid

Residual Diagnostics for Linear, Multilevel and Mixed-Effects Models

Description

This function performs residual diagnostics for linear models estimated by using the `lm()` function and for multilevel and linear mixed-effects models estimated by using the `lmer()` function from the **lme4** package to detect nonlinearity (partial residual or component-plus-residual plots), nonconstant error variance (predicted values vs. residuals plot), and non-normality of residuals (Q-Q plot and histogram with density plot).

Usage

```
check.resid(model, type = c("linear", "homo", "normal"),
            resid = c("unstand", "stand", "student"), plot = TRUE,
            point.shape = 21, point.fill = "gray80", point.size = 1,
            line1 = TRUE, line2 = TRUE, linetype1 = "solid",
            linetype2 = "dashed", linewidth1 = 1, linewidth2 = 1,
            line.col1 = "#0072B2", line.col2 = "#D55E00", bar.width = NULL,
            bar.n = 30, bar.col = "black", bar.fill = "gray95",
            strip.text.size = 11, label.size = 10, axis.text.size = 10,
            xlim = NULL, ylim = NULL, xbreaks = ggplot2::waiver(),
            ybreaks = ggplot2::waiver(), check = TRUE)
```

Arguments

<code>model</code>	a fitted model of class "lm", "lmerMod", or "lmerModLmerTest".
<code>type</code>	a character string specifying the type of the plot, i.e., "linear" for partial (component-plus-residual) plots, "homo" (default) for predicted values vs. residuals plot, and "normal" for Q-Q plot and histogram with a density plot. Note that partial residual or component-plus-residual plots are not available for models with interaction terms
<code>resid</code>	a character string specifying the type of residual used for the partial (component-plus-residual) plots or Q-Q plot and histogram, i.e., "unstand" for unstandardized residuals "stand" for standardized residuals, and "student" for studentized residuals. By default, studentized residuals are used for predicted values vs. residuals plot and unstandardized residuals are used for Q-Q plot and histogram. Note that studentized residuals are not available for multilevel and linear mixed-effects models when requesting Q-Q plots and histograms.

plot	logical: if TRUE (default), a plot is drawn.
point.shape	a numeric value for specifying the argument shape in the geom_point function.
point.fill	a character string or numeric value for specifying the argument fill the geom_point function.
point.size	a numeric value for specifying the argument size in the geom_point function.
line1	logical: if TRUE (default), regression line is drawn in the partial (component-plus-residual) plots, horizontal line is drawn in the predicted values vs. residuals plot, and t-distribution or normal distribution curve is drawn in the histogram.
line2	logical: if TRUE (default), Loess smooth line is drawn in the partial (component-plus-residual) plots, loess smooth lines are drawn in the predicted values vs. residuals plot, and density curve is drawn in the histogram.
linetype1	a character string or numeric value for specifying the argument linetype in the geom_smooth, geom_hline, or stat_function function.
linetype2	a character string or numeric value for specifying the argument linetype in the geom_smooth or geom_density function.
linewidth1	a numeric value for specifying the argument linewidth in the geom_smooth, geom_hline, or stat_function function.
linewidth2	a numeric value for specifying the argument linewidth in the geom_smooth or geom_density function.
line.col1	a character string or numeric value for specifying the argument color in the geom_smooth, geom_hline, or stat_function function.
line.col2	a character string or numeric value for specifying the argument color in the geom_smooth or geom_density function.
bar.width	a numeric value for specifying the argument bins in the geom_bar function.
bar.n	a numeric value for specifying the argument bins in the geom_bar function.
bar.col	a character string or numeric value for specifying the argument color in the geom_bar function.
bar.fill	a character string or numeric value for specifying the argument fill in the geom_bar function.
strip.text.size	a numeric value for specifying the argument size in the element_text function of the strip.text argument within the theme function.
label.size	a numeric value for specifying the argument size in the element_text function of the axis.title argument within the theme function.
axis.text.size	a numeric value for specifying the argument size in the element_text function of the axis.text argument within the theme function.
xlim	a numeric vector with two elements for specifying the argument limits in the scale_x_continuous function.
ylim	a numeric vector with two elements for specifying the argument limits in the scale_y_continuous function.
xbreaks	a numeric vector for specifying the argument breaks in the scale_x_continuous function.
ybreaks	a numeric vector for specifying the argument breaks in the scale_y_continuous function.
check	logical: if TRUE (default), argument specification is checked.

Details

Nonlinearity The violation of the assumption of linearity implies that the model cannot accurately capture the systematic pattern of the relationship between the outcome and predictor variables. In other words, the specified regression surface does not accurately represent the relationship between the conditional mean values of Y and the X s. That means the average error $E(\varepsilon)$ is not 0 at every point on the regression surface (Fox, 2015).

In multiple regression, plotting the outcome variable Y against each predictor variable X can be misleading because it does not reflect the partial relationship between Y and X (i.e., statistically controlling for the other X s), but rather the marginal relationship between Y and X (i.e., ignoring the other X s). Partial residual plots or component-plus-residual plots should be used to detect nonlinearity in multiple regression. The partial residual for the j th predictor variable is defined as

$$e_i^{(j)} = b_j X_{ij} + e_i$$

The linear component of the partial relationship between Y and X_j is added back to the least-squares residuals, which may include an unmodeled nonlinear component. Then, the partial residual $e_i^{(j)}$ is plotted against the predictor variable X_j . Nonlinearity may become apparent when a non-parametric regression smoother is applied.

By default, the function plots each predictor against the partial residuals, and draws the linear regression and the loess smooth line to the partial residual plots.

Nonconstant Error Variance The violation of the assumption of constant error variance, often referred to as heteroscedasticity, implies that the variance of the outcome variable around the regression surface is not the same at every point on the regression surface (Fox, 2015).

Plotting residuals against the outcome variable Y instead of the predicted values $\hat{Y}$ is not recommended because $Y = \hat{Y} + e$. Consequently, the linear correlation between the outcome variable Y and the residuals e is $\sqrt{1 - R^2}$ where R is the multiple correlation coefficient. In contrast, plotting residuals against the predicted values $\hat{Y}$ is much easier to examine for evidence of nonconstant error variance as the correlation between $\hat{Y}$ and e is 0. Note that the least-squares residuals generally have unequal variance $Var(e_i) = \sigma^2/(1 - h_i)$ where h is the leverage of observation i , even if errors have constant variance σ^2 . The studentized residuals e_i^* , however, have a constant variance under the assumption of the regression model. Residuals are studentized by dividing them by $\sigma_i^2(\sqrt{1 - h_i})$ where σ_i^2 is the estimate of σ^2 obtained after deleting the i th observation, and h_i is the leverage of observation i (Meuleman et al, 2015).

By default, the function plots the predicted values against the studentized residuals. It also draws a horizontal line at 0, a loess smooth lines for all residuals as well as separate loess smooth lines for positive and negative residuals.

Non-normality of Residuals Statistical inference under the violation of the assumption of normally distributed errors is approximately valid in all but small samples. However, the efficiency of least squares is not robust because the least-squares estimator is the most efficient and unbiased estimator only when the errors are normally distributed. For instance, when error distributions have heavy tails, the least-squares estimator becomes much less efficient compared to robust estimators. In addition, error distributions with heavy-tails result in outliers and compromise the interpretation of conditional means because the mean is not an accurate measure of central tendency in a highly skewed distribution. Moreover, a multimodal error

distribution suggests the omission of one or more discrete explanatory variables that naturally divide the data into groups (Fox, 2016).

By default, the function plots a Q-Q plot of the unstandardized residuals, and a histogram of the unstandardized residuals and a density plot. Note that studentized residuals follow a t -distribution with $n - k - 2$ degrees of freedom where n is the sample size and k is the number of predictors. However, the normal and t -distribution are nearly identical unless the sample size is small. Moreover, even if the model is correct, the studentized residuals are not an independent random sample from t_{n-k-2} . Residuals are correlated with each other depending on the configuration of the predictor values. The correlation is generally negligible unless the sample size is small.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>model</code>	model specified in <code>model</code>
<code>args</code>	specification of function arguments
<code>plotdat</code>	data frame used for the plot
<code>plot</code>	ggplot2 object for plotting the residuals

Note

This function uses a modified copy of the `partial()` and `calc_ranef()` function in the **remef** package by Sven Hohenstein and Reinhold Kliegl (2025) when requesting partial residual plots for linear mixed-effects models.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Fox, J. (2016). *Applied regression analysis and generalized linear models* (3rd ed.). Sage Publications, Inc.
- Hohenstein, S., & Kliegl, R. (2025). *remef: Remove Partial Effects*. R package version 1.0.7, <https://github.com/hohenstein/remef>
- Meuleman, B., Loosveldt, G., & Emonds, V. (2015). Regression analysis: Assumptions and diagnostics. In H. Best & C. Wolf (Eds.), *The SAGE handbook of regression analysis and causal inference* (pp. 83-110). Sage.

See Also

[check.collin](#), [check.outlier](#)

Examples

```

#-----
# Linear Model

# Estimate linear model
mod.lm <- lm(Ozone ~ Solar.R + Wind + Temp, data = airquality)

# Example 1a: Partial (component-plus-residual) plots
check.resid(mod.lm, type = "linear")

# Example 1b: Predicted values vs. residuals plot
check.resid(mod.lm, type = "homo")

# Example 1c: Q-Q plot and histogram with density plot
check.resid(mod.lm, type = "normal")

# Example 1d: Extract data and ggplot2 object
object <- check.resid(mod.lm, type = "linear", plot = FALSE)

# Data frame
object$plotdat

# ggplot object
object$plot

## Not run:
#-----
# Multilevel and Linear Mixed-Effects Model

# Estimate two-level mixed-effects model
mod.lmer <- lmer(Reaction ~ Days + (Days | Subject), data = sleepstudy)

# Example 2a: Partial (component-plus-residual) plots
check.resid(mod.lmer, type = "linear")

# Example 2b: Predicted values vs. residuals plot
check.resid(mod.lmer, type = "homo")

# Example 2c: Q-Q plot and histogram with density plot
check.resid(mod.lmer, type = "normal")

## End(Not run)

```

chr.color

*Colored and Styled Terminal Output Text***Description**

This function adds color and style to output texts on terminals that support 'ANSI' color and high-light codes that can be printed by using the cat function.

Usage

```
chr.color(x, color = c("black", "red", "green", "yellow", "blue", "violet",  
                      "cyan", "white", "gray1", "gray2", "gray3",  
                      "b.red", "b.green", "b.yellow", "b.blue", "b.violet",  
                      "b.cyan", "b.white"),  
         bg = c("none", "black", "red", "green", "yellow", "blue", "violet",  
               "cyan", "white"),  
         style = c("regular", "bold", "italic", "underline"), check = TRUE)
```

Arguments

x	a character vector.
color	a character string indicating the text color, e.g., red for red and b. red for bright red text.
bg	a character string indicating the background color of the text, e.g., red for red background.
style	a character vector indicating the font style, i.e., regular, (default) for regular text, bold for bold text, italic, for italic text, and underline for underline text. Note that font styles can be combined, e.g., style = c("bold", "italic") provides a bold and italic text.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a character vector.

Note

This function is based on functions provided in the **crayon** package by Gábor Csárdi.

Author(s)

Takuya Yanagida

References

Csárdi G (2022). *crayon: Colored Terminal Output*. R package version 1.5.2, <https://CRAN.R-project.org/package=crayon>

See Also

[chr.grep](#), [chr.grepl](#), [chr.gsub](#), [chr.omit](#), [chr.trim](#), [chr.trunc](#)

Examples

```
## Not run:

# Example 1:
cat(chr.color("Text in red.", color = "red"))

# Example 2:
cat(chr.color("Text in blue with green background.",
              color = "blue", bg = "yellow"))

# Example 3a:
cat(chr.color("Text in boldface.", style = "bold"))

# Example 3b:
cat(chr.color("Text in boldface and italic.", style = c("bold", "italic")))

## End(Not run)
```

chr.grep

Multiple Pattern Matching

Description

This function searches for matches to the character vector specified in `pattern` within each element of the character vector `x`.

Usage

```
chr.grep(pattern, x, ignore.case = FALSE, perl = FALSE, value = FALSE,
          fixed = FALSE, useBytes = FALSE, invert = FALSE, check = TRUE)

chr.grepl(pattern, x, ignore.case = FALSE, perl = FALSE, fixed = FALSE,
            useBytes = FALSE, check = TRUE)
```

Arguments

<code>pattern</code>	a character vector with character strings to be matched.
<code>x</code>	a character vector where matches are sought.
<code>ignore.case</code>	logical: if FALSE (default), the pattern matching is case sensitive and if TRUE, case is ignored during matching.
<code>perl</code>	logical: if TRUE Perl-compatible regexps are used.
<code>value</code>	logical: if FALSE (default), a vector containing the (integer) indices of the matches determined by <code>grep</code> is returned, and if TRUE, a vector containing the matching elements themselves is returned.
<code>fixed</code>	logical: if TRUE, <code>pattern</code> is a string to be matched as is. Overrides all conflicting arguments.

useBytes	logical: if TRUE, the matching is done byte-by-byte rather than character-by-character. See 'Details'.
invert	logical: if TRUE, function returns indices or values for elements that do not match.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a integer vector with the indices of the matches when value = FALSE, character vector containing the matching elements when value = TRUE, or a logical vector when using the chr.grep1 function.

Author(s)

Takuya Yanagida

References

Becker, R. A., Chambers, J. M. and Wilks, A. R. (1988) *The New S Language*. Wadsworth & Brooks/Cole

See Also

[chr.color](#), [chr.grep1](#), [chr.gsub](#), [chr.omit](#), [chr.trim](#), [chr.trunc](#)

Examples

```
chr.vector <- c("James", "Mary", "Michael", "Patricia", "Robert", "Jennifer")

# Example 1: Indices of matching elements
chr.grep(c("am", "er"), chr.vector)

# Example 2: Values of matching elements
chr.grep(c("am", "er"), chr.vector, value = TRUE)

# Example 3: Matching element?
chr.grep1(c("am", "er"), chr.vector)
```

chr.gsub

Multiple Pattern Matching And Replacements

Description

This function is a multiple global string replacement wrapper that allows access to multiple methods of specifying matches and replacements.

Usage

```
chr.gsub(pattern, replacement, x, recycle = FALSE, check = TRUE, ...)
```

Arguments

pattern	a character vector with character strings to be matched.
replacement	a character vector equal in length to pattern or of length one which are a replacement for matched patterns.
x	a character vector where matches and replacements are sought.
recycle	logical: if TRUE, replacement is recycled if lengths differ.
check	logical: if TRUE (default), argument specification is checked.
...	additional arguments to pass to the <code>regexr</code> or <code>sub</code> function.

Value

Return a character vector of the same length and with the same attributes as `x` (after possible coercion to character).

Note

This function was adapted from the `mgsub()` function in the **mgsub** package by Mark Ewing (2019).

Author(s)

Mark Ewing

References

Mark Ewing (2019). *mgsub: Safe, Multiple, Simultaneous String Substitution*. R package version 1.7.1. <https://CRAN.R-project.org/package=mgsub>

See Also

[chr.color](#), [chr.grep](#), [chr.grepl](#), [chr.omit](#), [chr.trim](#), [chr.trunc](#)

Examples

```
# Example 1: Replace 'the' and 'they' with 'a' and 'we'
chr.vector <- "they don't understand the value of what they seek."
chr.gsub(c("the", "they"), c("a", "we"), chr.vector)

# Example 2: Replace 'hey' and 'ho' with 'yo'
chr.vector <- c("hey ho, let's go!")
chr.gsub(c("hey", "ho"), "yo", chr.vector, recycle = TRUE)

# Example 3: Replace with regular expressions
chr.vector <- "Dopazamine is not the same as dopachloride or dopastriamine, yet is still fake."
chr.gsub(c("[Dd]opa([ ^]*?mine)", "fake"), c("Meta\\1", "real"), chr.vector)
```

chr.omit

*Omit Strings***Description**

This function omits user-specified values or strings from a numeric vector, character vector or factor.

Usage

```
chr.omit(x, omit = "", na.omit = FALSE, check = TRUE)
```

Arguments

x	a numeric vector, character vector or factor.
omit	a numeric vector or character vector indicating values or strings to be omitted from the vector x, the default setting is the empty strings "".
na.omit	logical: if TRUE, missing values (NA) are also omitted from the vector.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a numeric vector, character vector or factor with values or strings specified in `omit` omitted from the vector specified in `x`.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

See Also

[chr.color](#), [chr.grep](#), [chr.grepl](#), [chr.gsub](#), [chr.trim](#), [chr.trunc](#)

Examples

```
#-----
# Charater vector
x.chr <- c("a", "", "c", NA, "", "d", "e", NA)

# Example 1: Omit character string ""
chr.omit(x.chr)

# Example 2: Omit character string "" and missing values (NA)
chr.omit(x.chr, na.omit = TRUE)

# Example 3: Omit character string "c" and "e"
chr.omit(x.chr, omit = c("c", "e"))

# Example 4: Omit character string "c", "e", and missing values (NA)
```

```
chr.omit(x.chr, omit = c("c", "e"), na.omit = TRUE)
```

```
#-----
```

```
# Numeric vector
```

```
x.num <- c(1, 2, NA, 3, 4, 5, NA)
```

```
# Example 5: Omit values 2 and 4
```

```
chr.omit(x.num, omit = c(2, 4))
```

```
# Example 6: Omit values 2, 4, and missing values (NA)
```

```
chr.omit(x.num, omit = c(2, 4), na.omit = TRUE)
```

```
#-----
```

```
# Factor
```

```
x.factor <- factor(letters[1:10])
```

```
# Example 7: Omit factor levels "a", "c", "e", and "g"
```

```
chr.omit(x.factor, omit = c("a", "c", "e", "g"))
```

chr.trim

Trim Whitespace from String

Description

This function removes whitespace from start and/or end of a string

Usage

```
chr.trim(x, side = c("both", "left", "right"), check = TRUE)
```

Arguments

x	a character vector.
side	a character string indicating the side on which to remove whitespace, i.e., "both" (default), "left" or "right".
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a character vector with whitespaces removed from the vector specified in x.

Note

This function is based on the `str_trim()` function from the **stringr** package by Hadley Wickham.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Wickham, H. (2019). *stringr: Simple, consistent wrappers for common string operations*. R package version 1.4.0.

See Also

[chr.color](#), [chr.grep](#), [chr.grepl](#), [chr.gsub](#), [chr.omit](#), [chr.trunc](#)

Examples

```
x <- "  string  "

# Example 1: Remove whitespace at both sides
chr.trim(x)

# Example 2: Remove whitespace at the left side
chr.trim(x, side = "left")

# Example 3: Remove whitespace at the right side
chr.trim(x, side = "right")
```

chr.trunc	<i>Truncate a Character Vector to a Maximum Width</i>
-----------	---

Description

This function truncates a character vector, so that the number of characters of each element of the character vector is always less than or equal to the width specified in the argument width.

Usage

```
chr.trunc(x, width, side = c("right", "left", "center"), ellipsis = "...",
          check = TRUE)
```

Arguments

x	a character vector or factor. Note that factors are converted into a character vector.
width	a numeric value indicating the maximum width of the character strings in the vector. Note that the default setting switches to "." when width = 3, "." when width = 2, and "" when width = 1.
side	a character string indicating the location of the ellipsis, i.e. "right" (default) for the right side, "left" for the left side, and "center" for center of the character strings in the vector
ellipsis	a character string indicating the content of the ellipsis, i.e., "..." by default.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a truncated character vector.

Note

This function was adapted from the `str_trunc()` function in the **stringr** package by Hadley Wickham (2023).

Author(s)

Takuya Yanagida

References

Wickham H (2023). *stringr: Simple, Consistent Wrappers for Common String Operations*. R package version 1.5.1, <https://CRAN.R-project.org/package=stringr>

See Also

[chr.color](#), [chr.grep](#), [chr.grepl](#), [chr.gsub](#), [chr.omit](#), [chr.trim](#)

Examples

```
# Example 1: Truncate at the right side with a max. of 10 characters
chr.trunc(row.names(mtcars), width = 10)

# Example 2: Truncate at the left side with a max. of 10 characters
chr.trunc(row.names(mtcars), width = 10, side = "left")

# Example 3: Truncate without ellipses
chr.trunc(row.names(mtcars), width = 10, ellipsis = "")
```

ci.cor

(Bootstrap) Confidence Intervals for Correlation Coefficients

Description

This function computes and plots (1) Fisher z' confidence intervals for Pearson product-moment correlation coefficients (a) without non-normality adjustment, (1b) adjusted via sample joint moments method or (1c) adjusted via approximate distribution method (Bishara et al., 2018), (2) Spearman's rank-order correlation coefficients with (2a) Fieller et al. (1957) standard error, (2b) Bonett and Wright (2000) standard error, or (2c) rank-based inverse normal transformation, (3) Kendall's Tau-b, and (4) Kendall-Stuart's Tau-c correlation coefficients with Fieller et al. (1957) standard error, optionally by a grouping and/or split variable. The function also supports five types of bootstrap confidence intervals (e.g., bias-corrected (BC) percentile bootstrap or bias-corrected and accelerated (BCa) bootstrap confidence intervals) and plots the bootstrap samples with histograms and density curves. By default, the function computes Pearson product-moment correlation coefficients adjusted via approximate distribution method.

Usage

```
ci.cor(data, ...,
       method = c("pearson", "spearman", "kendall-b", "kendall-c"),
       adjust = c("none", "joint", "approx"),
       se = c("fisher", "fieller", "bonett", "rin"),
       sample = TRUE, seed = NULL, maxtol = 1e-05, nudge = 0.001,
       boot = c("none", "norm", "basic", "perc", "bc", "bca"), R = 1000,
       fisher = TRUE, alternative = c("two.sided", "less", "greater"),
       conf.level = 0.95, group = NULL, split = NULL, na.omit = FALSE, digits = 2,
       as.na = NULL, plot = c("none", "ci", "boot"), point.size = 2.5,
       point.shape = 19, errorbar.width = 0.3, dodge.width = 0.5, hist = TRUE,
       binwidth = NULL, bins = NULL, hist.alpha = 0.4, fill = "gray85", density = TRUE,
       density.col = "#0072B2", density.linewidth = 0.5, density.linetype = "solid",
       point = TRUE, point.col = "#CC79A7", point.linewidth = 0.6,
       point.linetype = "solid", ci = TRUE, ci.col = "black",
       ci.linewidth = 0.6, ci.linetype = "dashed", line = TRUE, intercept = 0,
       linetype = "solid", line.col = "gray65", xlab = NULL, ylab = NULL,
       xlim = NULL, ylim = NULL, xbreaks = ggplot2::waiver(), ybreaks = ggplot2::waiver(),
       axis.title.size = 11, axis.text.size = 10, strip.text.size = 11, title = NULL,
       subtitle = NULL, group.col = NULL, plot.margin = NA, legend.title = "",
       legend.position = c("right", "top", "left", "bottom", "none"),
       legend.box.margin = c(-10, 0, 0, 0), facet.ncol = NULL, facet.nrow = NULL,
       facet.scales = "free_y", filename = NULL, width = NA, height = NA,
       units = c("in", "cm", "mm", "px"), dpi = 600, write = NULL,
       append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a data frame with numeric variables, i.e., factors and character variables are excluded from data before conducting the analysis.
<code>...</code>	an expression indicating the variable names in data e.g., <code>ci.cor(x1, x2, data = dat)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the <code>df.subset</code> function.
<code>method</code>	a character string indicating which correlation coefficient is to be computed, i.e., "pearson" for Pearson product-moment correlation coefficient (default), "spearman" for Spearman's rank-order correlation coefficient, "kendall-b" for Kendall's Tau-b correlation coefficient, "kendall-c" for Kendall-Stuart's Tau-c correlation coefficient. Note that confidence intervals are only computed given at least 4 pairs of observations.
<code>adjust</code>	a character string specifying the non-normality adjustment method, i.e., "none" for the Fisher z' confidence interval for the Pearson product-moment correlation coefficient without non-normality adjustment, "joint" for the confidence interval with non-normality adjustment via sample joint moments, and "approx" (default) for the confidence interval with non-normality adjustment via approximate distribution by skewness and kurtosis. Note that this argument only applies to the Pearson product-moment correlation coefficient, i.e., <code>method = "pearson"</code>

se	a character string specifying the method for computing the standard error of the correlation coefficient, i.e., "fisher" for the Fisher z' confidence interval, "fieller" (default) for the confidence interval for Spearman's rank-order correlation coefficient based on approximate standard error by Fieller et al. (1957), "bonett" for the confidence interval based on approximate standard error by Bonett and Wright (2000), and "rin" for the confidence interval for Spearman's rank-order correlation coefficient based on rank-based inverse normal (RIN) transformation. Note that this argument only applies to Spearman's rank-order correlation coefficient, i.e., method = "spearman".
sample	logical: if TRUE (default), the univariate sample skewness and kurtosis is used when applying the approximate distribution method and reported in the result table, while the population skewness and kurtosis is used when sample = FALSE.
seed	a numeric value specifying the seed of the pseudo-random number generator when generating a random set of starting parameter value when the parameters led to a sum of squares greater than the maximum tolerance after optimization when applying the approximate distribution method (adjust = approx) when computing the confidence interval for the Pearson product-moment correlation coefficient, or the seeds of the pseudo-random numbers used when conducting bootstrapping.
maxtol	a numeric value indicating the tolerance for total squared error when applying the approximate distribution method (adjust = approx).
nudge	a numeric value indicating the nudge proportion of their original values by which sample skewness, kurtosis, and r are nudged towards 0 when applying the approximate distribution method (adjust = approx). are only computed given at least 10 pairs of observations.
boot	a character string specifying the type of bootstrap confidence intervals (CI), i.e., "none" (default) for not conducting bootstrapping, "norm" for the bias-corrected normal approximation bootstrap CI, "basic" for the basic bootstrap CI, "perc", for the percentile bootstrap CI "bc" (default) for the bias-corrected (BC) percentile bootstrap CI (without acceleration), and "bca" for the bias-corrected and accelerated (BCa) bootstrap CI.
R	a numeric value indicating the number of bootstrap replicates (default is 1000).
fisher	logical: if TRUE (default), Fisher z transformation is applied before computing the confidence intervals to reverse-transformed the limits of the interval using the inverse of the Fisher z transformation. Note that this argument applies only when boot is "norm" or "basic".
alternative	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
group	either a character string indicating the variable name of the grouping variable in ... or data, or a vector representing the grouping variable. The grouping variable is excluded from the data frame specified in data.
split	either a character string indicating the variable name of the split variable in ... or data, or a vector representing the split variable. The split variable is excluded from the data frame specified in data.

<code>na.omit</code>	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion); if FALSE (default), pairwise deletion is used.
<code>digits</code>	an integer value indicating the number of decimal places to be used.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>plot</code>	a character string indicating the type of the plot to display, i.e., "none" (default) for not displaying any plots, "ci" for displaying confidence intervals for the correlation coefficient, "boot" for displaying bootstrap samples with histograms and density curves when the argument "boot" is other than "none".
<code>point.size</code>	a numeric value indicating the size argument in the <code>geom_point</code> function for controlling the size of points when plotting confidence intervals (plot = "ci").
<code>point.shape</code>	a numeric value between 0 and 25 or a character string as plotting symbol indicating the shape argument in the <code>geom_point</code> function for controlling the symbols of points. when plotting confidence intervals (plot = "ci").
<code>errorbar.width</code>	a numeric value indicating the width argument in the <code>geom_errorbar</code> function for controlling the width of the whiskers in the <code>geom_errorbar</code> function when plotting confidence intervals (plot = "ci").
<code>dodge.width</code>	a numeric value indicating the width argument controlling the width of the geom elements to be dodged when specifying a grouping variable using the argument <code>group</code> when plotting confidence intervals (plot = "ci").
<code>hist</code>	logical: if TRUE (default), histograms are drawn when plotting bootstrap samples (plot = "boot").
<code>binwidth</code>	a numeric value or a function for specifying the binwidth argument in the <code>geom_histogram</code> function for controlling the width of the bins when plotting bootstrap samples (plot = "boot").
<code>bins</code>	a numeric value for specifying the bins argument in the <code>geom_histogram</code> function for controlling the number of bins when plotting bootstrap samples (plot = "boot").
<code>hist.alpha</code>	a numeric value between 0 and 1 for specifying the alpha argument in the <code>geom_histogram</code> function for controlling the opacity of the bars when plotting bootstrap samples (plot = "boot").
<code>fill</code>	a character string specifying the fill argument in the <code>geom_histogram</code> function controlling the fill aesthetic when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified <code>group = NULL</code> .
<code>density</code>	logical: if TRUE (default), density curves are drawn when plotting bootstrap samples (plot = "boot").
<code>density.col</code>	a character string specifying the color argument in the <code>geom_density</code> function controlling the color of the density curves when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified <code>group = NULL</code> .
<code>density.linewidth</code>	a numeric value specifying the linewidth argument in the <code>geom_density</code> function controlling the line width of the density curves when plotting bootstrap samples (plot = "boot").

density.linetype	a numeric value or character string specifying the linetype argument in the geom_density function controlling the line type of the density curves when plotting bootstrap samples (plot = "boot").
point	logical: if TRUE (default), vertical lines representing the point estimate of the correlation coefficients are drawn when plotting bootstrap samples (plot = "boot").
point.col	a character string specifying the color argument in the geom_vline function for controlling the color of the vertical line displaying the correlation coefficient when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
point.linewidth	a numeric value specifying the linewidth argument in the geom_vline function for controlling the line width of the vertical line displaying the correlation coefficient when plotting bootstrap samples (plot = "boot").
point.linetype	a numeric value or character string specifying the linetype argument in the geom_vline function controlling the line type of the vertical line displaying the correlation coefficient when plotting bootstrap samples (plot = "boot").
ci	logical: if TRUE (default), vertical lines representing the bootstrap confidence intervals of the correlation coefficient are drawn when plotting bootstrap samples (plot = "boot").
ci.col	character string specifying the color argument in the geom_vline function for controlling the color of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
ci.linewidth	a numeric value specifying the linewidth argument in the geom_vline function for controlling the line width of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (plot = "boot").
ci.linetype	a numeric value or character string specifying the linetype argument in the geom_vline function controlling the line type of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (plot = "boot").
line	logical: if TRUE (default), a horizontal line is drawn when plot = "ci" or a vertical line is drawn when plot = "boot"
intercept	a numeric value indicating the yintercept or xintercept argument in the geom_hline or geom_vline function controlling the position of the horizontal or vertical line when plot = "ci" and line = TRUE or when plot = "boot" and line = TRUE. By default, the horizontal or vertical line is drawn at 0.
linetype	a character string indicating the linetype argument in the geom_hline or geom_vline function controlling the line type of the horizontal or vertical line (default is linetype = "dashed").
line.col	a character string indicating the color argument in the geom_hline or geom_vline function for controlling the color of the horizontal or vertical line.
xlab	a character string indicating the name argument in the scale_x_continuous function for labeling the x-axis. The default setting is xlab = NULL when plot = "ci" and xlab = "Correlation Coefficient" when plot = "boot".

<code>ylab</code>	a character string indicating the name argument in the <code>scale_y_continuous</code> function for labeling the y-axis. The default setting is <code>ylab = "Correlation Coefficient"</code> when <code>plot = "ci"</code> and <code>ylab = "Probability Density, f(x)"</code> when <code>plot = "boot"</code> .
<code>xlim</code>	a numeric vector with two elements indicating the limits argument in the <code>scale_x_continuous</code> function for controlling the scale range of the x-axis. The default setting is <code>xlim = NULL</code> when <code>plot = "ci"</code> and <code>xlim = c(-1, 1)</code> when <code>plot = "boot"</code> .
<code>ylim</code>	a numeric vector with two elements indicating the limits argument in the <code>scale_y_continuous</code> function for controlling the scale range of the y-axis. The default setting is <code>ylim = c(-1, 1)</code> when <code>plot = "ci"</code> and <code>xlim = NULL</code> when <code>plot = "boot"</code> .
<code>xbreaks</code>	a numeric vector indicating the breaks argument in the <code>scale_x_continuous</code> function for controlling the x-axis breaks.
<code>ybreaks</code>	a numeric vector indicating the breaks argument in the <code>scale_y_continuous</code> function for controlling the y-axis breaks.
<code>axis.title.size</code>	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the axis title, i.e., <code>theme(axis.title = element_text(size = axis.text.size))</code> .
<code>axis.text.size</code>	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the axis text, i.e., <code>theme(axis.text = element_text(size = axis.text.size))</code> .
<code>strip.text.size</code>	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the strip text, i.e., <code>theme(strip.text = element_text(size = strip.text.size))</code> .
<code>title</code>	a character string indicating the title argument in the <code>labs</code> function for the subtitle of the plot.
<code>subtitle</code>	a character string indicating the subtitle argument in the <code>labs</code> function for the subtitle of the plot.
<code>group.col</code>	a character vector indicating the color argument in the <code>scale_color_manual</code> and <code>scale_fill_manual</code> functions when specifying a grouping variable using the argument <code>group</code> .
<code>plot.margin</code>	a numeric vector with four elements indicating the <code>plot.margin</code> argument in the <code>theme</code> function controlling the plot margins. The default setting is <code>c(5.5, 5.5, 5.5, 5.5)</code> , but switches to <code>c(5.5, 5.5, -2.5, 5.5)</code> when specifying a grouping variable using the argument <code>group</code> .
<code>legend.title</code>	a character string indicating the color argument in the <code>labs</code> function for specifying the legend title when specifying a grouping variable using the argument <code>group</code> .
<code>legend.position</code>	a character string indicating the <code>legend.position</code> in the <code>theme</code> argument for controlling the position of the legend function when specifying a grouping variable using the argument <code>group</code> . By default, the legend is placed at the bottom the plot.

<code>legend.box.margin</code>	a numeric vector with four elements indicating the <code>legend.box.margin</code> argument in the <code>theme</code> function for controlling the margins around the full legend area when specifying a grouping variable using the argument <code>group</code> .
<code>facet.ncol</code>	a numeric value indicating the <code>ncol</code> argument in the <code>facet_wrap</code> function for controlling the number of columns when specifying a split variable using the argument <code>split</code> .
<code>facet.nrow</code>	a numeric value indicating the <code>nrow</code> argument in the <code>facet_wrap</code> function for controlling the number of rows when specifying a split variable using the argument <code>split</code> .
<code>facet.scales</code>	a character string indicating the <code>scales</code> argument in the <code>facet_wrap</code> function for controlling the scales shared across facets, i.e., <code>"fixed"</code> , <code>"free_x"</code> , <code>"free_y"</code> (default), or <code>"free"</code> when specifying a split variable using the argument <code>split</code> .
<code>filename</code>	a character string indicating the <code>filename</code> argument including the file extension in the <code>ggsave</code> function. Note that one of <code>".eps"</code> , <code>".ps"</code> , <code>".tex"</code> , <code>".pdf"</code> (default), <code>".jpeg"</code> , <code>".tiff"</code> , <code>".png"</code> , <code>".bmp"</code> , <code>".svg"</code> or <code>".wmf"</code> needs to be specified as file extension in the <code>filename</code> argument. Note that plots can only be saved when <code>plot = "ci"</code> or <code>plot = "boot"</code> .
<code>width</code>	a numeric value indicating the <code>width</code> argument (default is the size of the current graphics device) in the <code>ggsave</code> function.
<code>height</code>	a numeric value indicating the <code>height</code> argument (default is the size of the current graphics device) in the <code>ggsave</code> function.
<code>units</code>	a character string indicating the <code>units</code> argument (default is <code>in</code>) in the <code>ggsave</code> function.
<code>dpi</code>	a numeric value indicating the <code>dpi</code> argument (default is 600) in the <code>ggsave</code> function.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension <code>".txt"</code> (e.g., <code>"Output.txt"</code>) or Excel file with file extension <code>".xlsx"</code> (e.g., <code>"Output.xlsx"</code>). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if <code>TRUE</code> (default), output will be appended to an existing text file with extension <code>".txt"</code> specified in <code>write</code> , if <code>FALSE</code> existing text file will be overwritten.
<code>check</code>	logical: if <code>TRUE</code> (default), argument specification is checked.
<code>output</code>	logical: if <code>TRUE</code> (default), output is shown on the console.

Details

Pearson Product-Moment Correlation Coefficient The Fisher z' confidence interval method for the Pearson product-moment correlation coefficient is based on the assumption that X and Y have a bivariate normal distribution in the population. Non-normality resulting from either high kurtosis or high absolute skewness can distort the Fisher z' confidence interval that produces a coverage rate that does not equal the one intended. The distortion is largest when population correlation is large and both variables X and Y were non-normal (Bishara et al., 2017). Note that increasing sample size improves coverage only when the population correlation is zero, while increasing the sample size worsens coverage with a non-zero population

correlation (Bishara & Hittner, 2017). The `ci.cor` function computes the Fisher z' confidence interval without non-normality adjustment (`adjust = "none"`), with non-normality adjustment via sample joint moments (`adjust = "joint"`), or with non-normality adjustment via approximate distribution (`adjust = "approx"`):

- *Fisher z' confidence interval method* uses the r -to- z' transformation for the correlation coefficient r :

$$z' = 0.5 \cdot \ln \left(\frac{1+r}{1-r} \right)$$

The sampling distribution of z is approximately normal with a standard error of approximately

$$\sigma'_z = \sqrt{\frac{1}{(n-3)}}$$

The two-sided 95% confidence interval is defined as

$$z' \pm 1.96 \cdot \sigma_{z'}$$

These confidence interval bounds are transformed back to the scale of r :

$$r = \frac{\exp(2z') - 1}{\exp(2z') + 1}$$

The resulting confidence interval of the correlation coefficient is an approximation and is only accurate when X and Y have a bivariate normal distribution in the population or when the population correlation is zero.

- The *Joint Moments Method* multiplies the asymptotic variance of z' by τ_f^2 (Hawkins, 1989):

$$\tau_f^2 = \frac{(\mu_{40} + 2\mu_{22} + \mu_{04})\rho^2 - 4(\mu_{31} + \mu_{13})\rho + 4\mu_{22}}{4(1 - \rho^2)^2}$$

where μ_{jk} represents a population joint moment defined as

$$\mu_{jk} = E[X^j Y^k]$$

where X and Y are assumed to be standardized ($\mu_{10} = \mu_{01} = 0$, $\mu_{20} = \mu_{02} = 1$). The standard error of z' can then be approximated as $\tilde{\sigma}_{z'}$:

$$\tilde{\sigma}_{z'} = \tau_f \sqrt{\frac{1}{n-3}}$$

The corresponding sample moments, m_{jk} can be used to estimate τ_f^2 :

$$\hat{\mu}_{jk} = m_{jk} = \frac{1}{n} \sum_{i=1}^n (x_i^j y_i^k)$$

However, the higher-order sample joint moments may be unstable estimators of their population counterparts unless the sample size is extremely large. Thus, this estimate of τ_f^2 may be inaccurate, leading to inaccurate confidence intervals.

- The *Approximate Distribution Method* estimates an approximate distribution that the sample appears to be drawn from to analytically solve for τ_f^2 based on that distribution's parameters. The `ci.cor` function uses a third-order polynomial family allowing estimation of distribution parameters using marginal skewness and kurtosis that are estimated using the marginal sample skewness and kurtosis statistics (Bishara et al., 2018).

Bishara et al. (2018) conducted two Monte Carlo simulations that showed that the approximate distribution method was effective in dealing with violations of the bivariate normality assumption for a wide range of sample sizes, while the joint moments method was effective mainly when the sample size was extremely large, in the thousands. However, the third-order polynomial family used for the approximate distribution method cannot deal with absolute skewness above 4.4 or kurtosis above 43.4. Note that the approximate distribution method is accurate even when the bivariate normality assumption is satisfied, while the sample joint moments method sometimes fails to achieve the intended coverage even when the bivariate normality was satisfied.

Spearman's Rank-Order Correlation Coefficient The confidence interval for Spearman's rank-order correlation coefficient is based on the Fisher's z method (`se = "fisher"`), Fieller et al. (1957) approximate standard error (`se = "fieller"`, default), Bonett and Wright (2000) approximate standard error (`se = "bonett"`) or rank-based inverse normal (RIN) transformation (`se = "rin"`):

- *Fisher's z Standard Error*

$$\sqrt{\frac{1}{(n-3)}}$$

- *Fieller et al. (1957) Approximate Standard Error*

$$\sqrt{\frac{1.06}{(n-3)}}$$

Note that this approximation for the standard error is recommended for $n > 10$ and $|rs| < 0.8$.

- *Bonett and Wright (2000) Approximate Standard Error*

$$\sqrt{\frac{1 + \frac{\hat{\theta}^2}{2}}{(n-3)}}$$

where $\hat{\theta}$ is the point estimate of the Spearman's rank-order correlation coefficient. Note that this approximation for the standard error is recommended for $|\tau| \leq 0.9$.

- *Rin Transformation* involves three steps. First, the variable is converted to ranks. Second, the ranks are converted to a 0-to-1 scale using a linear function. Third, this distribution is transformed via the inverse of the normal cumulative distribution function (i.e., via probit transformation). The result is an approximately normal distribution regardless of the original shape of the data, so long as ties are infrequent and n is not too small.

Kendall's Tau-b and Tau-c Correlation Coefficient The confidence interval for Kendall's Tau-b and Tau-c correlation coefficient is based on the approximate standard error by Fieller et al. (1957):

$$\sigma'_z = \sqrt{\frac{0.437}{(n-4)}}$$

Note that this approximation for the standard error is recommended for $n > 10$ and $|\tau| < 0.8$.

Bootstrap Confidence Intervals The `ci.cor` function supports bootstrap confidence intervals (CI) for the correlation coefficient by changing the default setting `boot = "none"` to request one of five different types of bootstrap CI (see Efron & Tibshirani, 1993; Davidson & Hinkley, 1997):

- "norm": The bias-corrected normal approximation bootstrap CI relies on the normal distribution based on the standard deviation of the bootstrap samples $\hat{SE}^*$. The function corrects for the bootstrap bias, i.e., difference between the bootstrap estimate $\hat{\theta}^*$ and the sample statistic $\hat{\theta}$ centering the interval at $2\hat{\theta} - \hat{\theta}^*$. The BC normal CI of intended coverage of $1 - 2(\alpha/2)$ is given by

$$Normal : (\hat{\theta}_{low}, \hat{\theta}_{upp} = \hat{\theta} - (\hat{\theta}^* - \hat{\theta}) + z^{\alpha/2} \cdot \hat{SE}^*, \hat{\theta} - (\hat{\theta}^* - \hat{\theta}) + z^{1-\alpha/2} \cdot \hat{SE}^*)$$

where $z^{\alpha/2}$ and $z^{1-\alpha/2}$ denotes the α and the $1 - \alpha$ quantile from the standard normal distribution.

- "basic": The basic bootstrap (aka reverse bootstrap percentile) CI is based on the distribution of $\hat{\delta} = \hat{\theta} - \theta$ which is approximated with the bootstrap distribution of $\hat{\delta}^* = \hat{\theta}^* - \hat{\theta}$.

$$Basic : (\hat{\theta}_{low}, \hat{\theta}_{upp} = \hat{\theta} - \hat{\delta}^{*(1-\alpha/2)}, \hat{\theta} - \hat{\delta}^{*\alpha/2} = 2\hat{\theta} - \hat{\theta}^{*(1-\alpha/2)}, 2\hat{\theta} - \hat{\theta}^{*(\alpha/2)})$$

- "perc": The percentile bootstrap CI is computed by ordering the bootstrap estimates $\hat{\theta}_1^*, \dots, \hat{\theta}_B^*$ to determine the $(100(\alpha)/2)$ th and $(100(1 - \alpha)/2)$ th empirical percentile with intended coverage of $1 - 2(\alpha/2)$:

$$Percentile : (\hat{\theta}_{low}, \hat{\theta}_{upp} = \hat{\theta}^{*(1-\alpha/2)}, \hat{\theta}^{*(\alpha/2)})$$

- "bc" (default): The bias-corrected (BC) percentile bootstrap CI corrects the percentile bootstrap CI for median bias of $\hat{\theta}^*$, i.e., the discrepancy between the median of $\hat{\theta}^*$ and $\hat{\theta}$ in normal units. The bias correction $\hat{z}_0$ is obtained from the proportion of bootstrap replications less than the sample estimate $\hat{\theta}$:

$$\hat{z}_0 = \Phi^{-1} \left(\frac{\#\hat{\theta}_b^* < \hat{\theta}}{B} \right)$$

where $\Phi^{-1}(\cdot)$ represents the inverse function of the standard normal cumulative distribution function and B is the number of bootstrap replications. The BC percentile CI of intended coverage of $1 - 2(\alpha/2)$ is given by

$$BC : (\hat{\theta}_{low}, \hat{\theta}_{upp} = \hat{\theta}^{*(\alpha_1)}, \hat{\theta}^{*(\alpha_2)})$$

where

$$\alpha_1 = \Phi(2\hat{z}_0 + z^{\alpha/2})$$

$$\alpha_2 = \Phi(2\hat{z}_0 + z^{1-\alpha/2})$$

where $\Phi(\cdot)$ represents the standard normal cumulative distribution function and $z^{\alpha/2}$ is the $100(\alpha/2)$ percentile of a standard normal distribution.

- "bca": The bias-corrected and accelerated (BCa) bootstrap CI corrects the percentile bootstrap CI for median bias $\hat{z}_0$ and for acceleration or skewness $\hat{a}$, i.e., the rate of change of the standard error of $\hat{\theta}$ with respect to the true parameter value θ on a normalized scale. The standard normal approximation $\hat{\theta} \sim N(\theta, SE^2)$ assumes that the standard error of $\hat{\theta}$ is the same for all θ . The acceleration constant $\hat{a}$ corrects for this unrealistic assumption and can be computed by using jackknife resampling:

$$\hat{a} = \frac{\sum_{i=1}^n (\hat{\theta}_{(\cdot)} - \hat{\theta}_{(i)})^3}{6\{\sum_{i=1}^n (\hat{\theta}_{(\cdot)} - \hat{\theta}_{(i)})^2\}^{3/2}}$$

where $\hat{\theta}_{(i)}$ is the sample estimate with the i th observation deleted and $\hat{\theta}_{(\cdot)} = \sum_{i=1}^n \frac{\hat{\theta}_{(i)}}{n}$. Note that the function uses infinitesimal jackknife instead of regular leave-one-out jackknife that down-weights each observation by an infinitesimal amount of $\frac{0.001}{n}$ instead of removing observations. The BCa percentile CI of intended coverage of $1 - 2(\alpha/2)$ is given by

$$BCa : (\hat{\theta}_{low}, \hat{\theta}_{upp} = \hat{\theta}^{*(\alpha_1)}, \hat{\theta}^{*(\alpha_2)})$$

where

$$\alpha_1 = \Phi \left(\hat{z}_0 + \frac{\hat{z}_0 + z^{\alpha/2}}{1 - \hat{a}(\hat{z}_0 + z^{\alpha/2})} \right)$$

$$\alpha_2 = \Phi \left(\hat{z}_0 + \frac{\hat{z}_0 + z^{1-\alpha/2}}{1 - \hat{a}(\hat{z}_0 + z^{1-\alpha/2})} \right)$$

Note that Fisher transformation is applied before computing the confidence intervals to reverse-transform the limits of the interval using the inverse of the Fisher z transformation (fisher = TRUE) when specifying "norm" or "basic" for the argument boot. In addition, interpolation on the normal quantile scale is applied for "basic", "perc", "bc", and "bca" when a non-integer order statistic is required (see equation 5.8 in Davison & Hinkley, 1997).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	list with the input specified in <code>...</code> , <code>data</code> , <code>group</code> , and <code>split</code>
<code>args</code>	specification of function arguments
<code>boot</code>	data frame with bootstrap replicates of the correlation coefficient when bootstrapping was requested
<code>plot</code>	<code>ggplot2</code> object for plotting the results
<code>result</code>	result table

Note

This function is based on a modified copy of the functions provided in the supporting information in Bishara et al. (2018) for the sample joint moments method and approximate distribution method, functions provided in the supplementary materials in Bishara and Hittner (2017) for Fieller et al. (1957) and Bonett and Wright (2000) correction, and a function provided by Thom Baguley (2024) for the rank-based inverse normal (RIN) transformation. Bootstrap confidence intervals are computed using the R package *boot* by Angelo Canty and Brian Ripley (2024).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Baguley, T. (2024). *CI for Spearman's rho*. *seriousstats*. <https://rpubs.com/seriousstats/616206>
- Bonett, D. G., & Wright, T. A. (2000). Sample size requirements for estimating Pearson, Kendall and Spearman correlations. *Psychometrika*, 65, 23-28. <https://doi.org/10.1007/BF02294183>
- Bishara, A. J., & Hittner, J. B. (2012). Testing the significance of a correlation with nonnormal data: Comparison of Pearson, Spearman, transformation, and resampling approaches. *Psychological Methods*, 17(3), 399-417. <https://doi.org/10.1037/a0028087>
- Bishara, A. J., & Hittner, J.B. (2017). Confidence intervals for correlations when data are not normal. *Behavior Research Methods*, 49, 294-309. <https://doi.org/10.3758/s13428-016-0702-8>
- Bishara, A. J., Li, J., & Nash, T. (2018). Asymptotic confidence intervals for the Pearson correlation via skewness and kurtosis. *British Journal of Mathematical and Statistical Psychology*, 71(1), 167-185. <https://doi.org/10.1111/bmsp.12113>
- Brown, M. B., & Benedetti, J. K. (1977). Sampling behavior of tests for correlation in two-way contingency tables. *Journal of the American Statistical Association*, 72 (358), 309-315. <https://doi.org/10.1080/01621459.1977.1>
- Canty, A., & Ripley, B. (2024). *boot: Bootstrap R (S-Plus) Functions*. R package version 1.3-31.
- Davison, A. C., & Hinkley, D. V. (1997). *Bootstrap methods and their application*. Cambridge University Press.
- Efron, B., & Tibshirani, R. (1993). *An introduction to the bootstrap*. Chapman & Hall.
- Fieller, E. C., Hartley, H. O., & Pearson, E. S. (1957). Tests for rank correlation coefficients: I. *Biometrika*, 44, 470-481. <https://doi.org/10.2307/2332878>
- Fisher, R. A. (1921). On the "Probable Error" of a Coefficient of Correlation Deduced from a Small Sample. *Metron*, 1, 3-32.
- Hawkins, D. L. (1989). Using U statistics to derive the asymptotic distribution of Fisher's Z statistic. *American Statistician*, 43, 235-237. <https://doi.org/10.2307/2685369>
- Hollander, M., Wolfe, D. A., & Chicken, E. (2015). *Nonparametric statistical methods*. Wiley.

See Also

[cor.matrix](#), [ci.mean](#), [ci.mean.diff](#), [ci.prop](#), [ci.var](#), [ci.sd](#)

Examples

```

#-----
# Pearson product-moment correlation coefficient

# Example 1a: Approximate distribution method
ci.cor(mtcars, mpg, drat, qsec)

# Alternative specification without using the '...' argument
ci.cor(mtcars[, c("mpg", "drat", "qsec")])

# Example 1b: Joint moments method
ci.cor(mtcars, mpg, drat, qsec, adjust = "joint")

#-----
# Spearman's rank-order correlation coefficient

# Example 2a: Fieller et al. (1957) approximate standard error
ci.cor(mtcars, mpg, drat, qsec, method = "spearman")

# Example 2b: Bonett and Wright (2000) approximate standard error
ci.cor(mtcars, mpg, drat, qsec, method = "spearman", se = "bonett")

# Example 2c: Rank-based inverse normal (RIN) transformation
ci.cor(mtcars, mpg, drat, qsec, method = "spearman", se = "rin")

#-----
# Kendall's Tau

# Example 3a: Kendall's Tau-b
ci.cor(mtcars, mpg, drat, qsec, method = "kendall-b")

# Example 3b: Kendall's Tau-c
ci.cor(mtcars, mpg, drat, qsec, method = "kendall-c")

## Not run:
#-----
# Bootstrap Confidence Interval (CI)

# Example 4a: Bias-corrected (BC) percentile bootstrap CI
ci.cor(mtcars, mpg, drat, qsec, boot = "bc")

# Example 4b: Bias-corrected and accelerated (BCa) bootstrap CI,
# 5000 bootstrap replications, set seed of the pseudo-random number generator
ci.cor(mtcars, mpg, drat, qsec, boot = "bca", R = 5000, seed = 123)

#-----
# Grouping and Split Variable

# Example 5a: Grouping variable
ci.cor(mtcars, mpg, drat, qsec, group = "vs")

# Alternative specification without using the argument '...'

```

```

ci.cor(mtcars[, c("mpg", "drat", "qsec")], group = mtcars$vs)

# Example 5b: Split variable
ci.cor(mtcars, mpg, drat, qsec, split = "am")

# Alternative specification without using the argument '...'
ci.cor(mtcars[, c("mpg", "drat", "qsec")], split = mtcars$am)

# Example 5c: Grouping and split variable
ci.cor(mtcars, mpg, drat, qsec, group = "vs", split = "am")

# Alternative specification without using the argument '...'
ci.cor(mtcars[, c("mpg", "drat", "qsec")], group = mtcars$vs, split = mtcars$am)

#-----
# Write Output

# Example 6a: Text file
ci.cor(mtcars, mpg, drat, qsec, write = "CI_Cor_Text.txt")

# Example 6b: Excel file
ci.cor(mtcars, mpg, drat, qsec, write = "CI_Cor_Excel.xlsx")

#-----
# Plot Confidence Intervals

# Example 7a: Pearson product-moment correlation coefficient
ci.cor(mtcars, mpg, drat, qsec, plot = "ci")

# Example 7b: Grouping variable
ci.cor(mtcars, mpg, drat, qsec, group = "vs", plot = "ci")

# Example 7c: Split variable
ci.cor(mtcars, mpg, drat, qsec, split = "am", plot = "ci")

# Example 7d: Save plot as PDF file
ci.cor(mtcars, mpg, drat, qsec, plot = "ci", saveplot = "CI_Cor.pdf",
       width = 8, height = 6)

# Example 7e: Save plot as PNG file
ci.cor(mtcars, mpg, drat, qsec, plot = "ci", saveplot = "CI_Cor.png",
       width = 8, height = 6)

#-----
# Plot Bootstrap Samples

# Example 8a: Pearson product-moment correlation coefficient
ci.cor(mtcars, mpg, drat, qsec, boot = "bc", plot = "boot")

# Example 8b: Grouping variable
ci.cor(mtcars, mpg, drat, qsec, group = "vs", boot = "bc", plot = "boot")

# Example 8c: Split variable

```

```

ci.cor(mtcars, mpg, drat, qsec, split = "am", boot = "bc", plot = "boot")

# Example 8d: Save plot as PDF file
ci.cor(mpg, drat, qsec, data = mtcars, boot = "bc", plot = "boot",
       saveplot = "CI_Cor_Boot.pdf", width = 14, height = 9)

# Example 8e: Save plot as PNG file
ci.cor(mtcars, mpg, drat, qsec, boot = "bc", plot = "boot",
       saveplot = "CI_Cor_Boot.png", width = 14, height = 9)
## End(Not run)

```

ci.mean

(Bootstrap) Confidence Intervals for Arithmetic Means and Medians

Description

The function `ci.mean` computes and plots confidence intervals for arithmetic means with known or unknown population standard deviation or population variance and the function `ci.median` computes confidence intervals for medians, optionally by a grouping and/or split variable. These functions also supports six types of bootstrap confidence intervals (e.g., bias-corrected (BC) percentile bootstrap or bias-corrected and accelerated (BCa) bootstrap confidence intervals) and plots the bootstrap samples with histograms and density curves.

Usage

```

ci.mean(data, ..., sigma = NULL, sigma2 = NULL, adjust = FALSE,
        boot = c("none", "norm", "basic", "stud", "perc", "bc", "bca"),
        R = 1000, seed = NULL, sample = TRUE,
        alternative = c("two.sided", "less", "greater"),
        conf.level = 0.95, group = NULL, split = NULL, sort.var = FALSE,
        na.omit = FALSE, digits = 2, as.na = NULL,
        plot = c("none", "ci", "boot"), point.size = 2.5, point.shape = 19,
        errorbar.width = 0.3, dodge.width = 0.5, hist = TRUE,
        binwidth = NULL, bins = NULL, hist.alpha = 0.4, fill = "gray85", density = TRUE,
        density.col = "#0072B2", density.linewidth = 0.5, density.linetype = "solid",
        point = TRUE, point.col = "#CC79A7", point.linewidth = 0.6,
        point.linetype = "solid", ci = TRUE, ci.col = "black",
        ci.linewidth = 0.6, ci.linetype = "dashed", line = FALSE, intercept = 0,
        linetype = "solid", line.col = "gray65", xlab = NULL, ylab = NULL,
        xlim = NULL, ylim = NULL, xbreaks = ggplot2::waiver(),
        ybreaks = ggplot2::waiver(), axis.title.size = 11, axis.text.size = 10,
        strip.text.size = 11, title = NULL, subtitle = NULL, group.col = NULL,
        plot.margin = NA, legend.title = "",
        legend.position = c("right", "top", "left", "bottom", "none"),
        legend.box.margin = c(-10, 0, 0, 0), facet.ncol = NULL, facet.nrow = NULL,
        facet.scales = "free", filename = NULL, width = NA, height = NA,
        units = c("in", "cm", "mm", "px"), dpi = 600, write = NULL,
        append = TRUE, check = TRUE, output = TRUE)

```

```

ci.median(data, ..., boot = c("none", "norm", "basic", "stud", "perc", "bc", "bca"),
  R = 1000, seed = NULL, sample = TRUE,
  alternative = c("two.sided", "less", "greater"),
  conf.level = 0.95, group = NULL, split = NULL, sort.var = FALSE,
  na.omit = FALSE, digits = 2, as.na = NULL, plot = c("none", "ci", "boot"),
  point.size = 2.5, point.shape = 19, errorbar.width = 0.3, dodge.width = 0.5,
  hist = TRUE, binwidth = NULL, bins = NULL, hist.alpha = 0.4, fill = "gray85",
  density = TRUE, density.col = "#0072B2", density.linewidth = 0.5,
  density.linetype = "solid", point = TRUE, point.col = "#CC79A7",
  point.linewidth = 0.6, point.linetype = "solid", ci = TRUE, ci.col = "black",
  ci.linewidth = 0.6, ci.linetype = "dashed", line = FALSE, intercept = 0,
  linetype = "solid", line.col = "gray65", xlab = NULL, ylab = NULL,
  xlim = NULL, ylim = NULL, xbreaks = ggplot2::waiver(),
  ybreaks = ggplot2::waiver(), axis.title.size = 11, axis.text.size = 10,
  strip.text.size = 11, title = NULL, subtitle = NULL, group.col = NULL,
  plot.margin = NA, legend.title = "",
  legend.position = c("right", "top", "left", "bottom", "none"),
  legend.box.margin = c(-10, 0, 0, 0), facet.ncol = NULL, facet.nrow = NULL,
  facet.scales = "free", filename = NULL, width = NA, height = NA,
  units = c("in", "cm", "mm", "px"), dpi = 600, write = NULL, append = TRUE,
  check = TRUE, output = TRUE)

```

Arguments

data	a numeric vector or data frame with numeric variables, i.e., factors and character variables are excluded from data before conducting the analysis.
...	an expression indicating the variable names in data e.g., <code>ci.mean(x1, x2, data = dat)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
sigma	a numeric vector indicating the population standard deviation when computing confidence intervals for the arithmetic mean with known standard deviation Note that either argument <code>sigma</code> or argument <code>sigma2</code> is specified and it is only possible to specify one value for the argument <code>sigma</code> even though multiple variables are specified in data.
sigma2	a numeric vector indicating the population variance when computing confidence intervals for the arithmetic mean with known variance. Note that either argument <code>sigma</code> or argument <code>sigma2</code> is specified and it is only possible to specify one value for the argument <code>sigma2</code> even though multiple variables are specified in data.
adjust	logical: if TRUE, difference-adjustment for the confidence intervals for the arithmetic mean (Baguley, 2012) is applied.
boot	a character string specifying the type of bootstrap confidence intervals (CI), i.e., "none" (default) for not conducting bootstrapping, "norm" for the bias-corrected normal approximation bootstrap CI, "basic" for the basic bootstrap CI, "stud" for the studentized bootstrap CI, "perc", for the percentile bootstrap

	CI "bc" for the bias-corrected (BC) percentile bootstrap CI (without acceleration), and "bca" for the bias-corrected and accelerated (BCa) bootstrap CI, see 'Details' in the ci.cor function.
R	a numeric value indicating the number of bootstrap replicates (default is 1000).
seed	a numeric value specifying seeds of the pseudo-random numbers used in the bootstrap algorithm when conducting bootstrapping.
sample	logical: if TRUE (default), the univariate sample skewness and kurtosis is computed, while the population skewness and kurtosis is computed when sample = FALSE.
alternative	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
group	either a character string indicating the variable name of the grouping variable in data, or a vector representing the grouping variable. The grouping variable is excluded from the data frame specified in data. Note that a grouping variable can only be used when computing confidence intervals with unknown population standard deviation and population variance.
split	either a character string indicating the variable name of the split variable in data, or a vector representing the split variable. The split variable is excluded from the data frame specified in data. Note that a grouping variable can only be used when computing confidence intervals with unknown population standard deviation and population variance.
sort.var	logical: if TRUE, output table is sorted by variables when specifying group.
na.omit	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion) when specifying more than one outcome variable.
digits	an integer value indicating the number of decimal places to be used.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that as.na() function is only applied to data, but not to group or split.
plot	a character string indicating the type of the plot to display, i.e., "none" (default) for not displaying any plots, "ci" for displaying confidence intervals for the arithmetic mean or median, "boot" for displaying bootstrap samples with histograms and density curves when the argument "boot" is other than "none".
point.size	a numeric value indicating the size argument in the geom_point function for controlling the size of points when plotting confidence intervals (plot = "ci").
point.shape	a numeric value between 0 and 25 or a character string as plotting symbol indicating the shape argument in the geom_point function for controlling the symbols of points when plotting confidence intervals (plot = "ci").
errorbar.width	a numeric value indicating the width argument in the geom_errorbar function for controlling the width of the whiskers in the geom_errorbar function when plotting confidence intervals (plot = "ci").
dodge.width	a numeric value indicating the width argument controlling the width of the geom elements to be dodged when specifying a grouping variable using the argument group and plotting confidence intervals (plot = "ci").

hist	logical: if TRUE (default), histograms are drawn when plotting bootstrap samples (plot = "boot").
binwidth	a numeric value or a function for specifying the binwidth argument in the geom_histogram function for controlling the width of the bins when plotting bootstrap samples (plot = "boot").
bins	a numeric value for specifying the bins argument in the geom_histogram function for controlling the number of bins when plotting bootstrap samples (plot = "boot").
hist.alpha	a numeric value between 0 and 1 for specifying the alpha argument in the geom_histogram function for controlling the opacity of the bars when plotting bootstrap samples (plot = "boot").
fill	a character string specifying the fill argument in the geom_histogram function controlling the fill aesthetic when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
density	logical: if TRUE (default), density curves are drawn when plotting bootstrap samples (plot = "boot").
density.col	a character string specifying the color argument in the geom_density function controlling the color of the density curves when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
density.linewidth	a numeric value specifying the linewidth argument in the geom_density function controlling the line width of the density curves when plotting bootstrap samples (plot = "boot").
density.linetype	a numeric value or character string specifying the linetype argument in the geom_density function controlling the line type of the density curves when plotting bootstrap samples (plot = "boot").
point	logical: if TRUE (default), vertical lines representing the point estimate of the arithmetic mean or median are drawn when plotting bootstrap samples (plot = "boot").
point.col	a character string specifying the color argument in the geom_vline function for controlling the color of the vertical line displaying the arithmetic mean or median when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
point.linewidth	a numeric value specifying the linewidth argument in the geom_vline function for controlling the line width of the vertical line displaying the arithmetic mean or median when plotting bootstrap samples (plot = "boot").
point.linetype	a numeric value or character string specifying the linetype argument in the geom_vline function controlling the line type of the vertical line displaying the arithmetic mean or median when plotting bootstrap samples (plot = "boot").
ci	logical: if TRUE (default), vertical lines representing the bootstrap confidence intervals of the arithmetic mean or median are drawn when plotting bootstrap samples (plot = "boot").

ci.col	character string specifying the color argument in the <code>geom_vline</code> function for controlling the color of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (<code>plot = "boot"</code>). Note that this argument applied only when no grouping variable was specified <code>group = NULL</code> .
ci.linewidth	a numeric value specifying the <code>linewidth</code> argument in the <code>geom_vline</code> function for controlling the line width of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (<code>plot = "boot"</code>).
ci.linetype	a numeric value or character string specifying the <code>linetype</code> argument in the <code>geom_vline</code> function controlling the line type of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (<code>plot = "boot"</code>).
line	logical: if TRUE, a horizontal line is drawn when <code>plot = "ci"</code> or a vertical line is drawn when <code>plot = "boot"</code>
intercept	a numeric value indicating the <code>yintercept</code> or <code>xintercept</code> argument in the <code>geom_hline</code> or <code>geom_vline</code> function controlling the position of the horizontal or vertical line when <code>plot = "ci"</code> and <code>line = TRUE</code> or when <code>plot = "boot"</code> and <code>line = TRUE</code> . By default, the horizontal or vertical line is drawn at 0.
linetype	a character string indicating the <code>linetype</code> argument in the <code>geom_hline</code> or <code>geom_vline</code> function controlling the line type of the horizontal or vertical line (default is <code>linetype = "dashed"</code>).
line.col	a character string indicating the color argument in the <code>geom_hline</code> or <code>geom_vline</code> function for controlling the color of the horizontal or vertical line.
xlab	a character string indicating the name argument in the <code>scale_x_continuous</code> function for labeling the x-axis. The default setting is <code>xlab = NULL</code> when <code>plot = "ci"</code> and <code>xlab = "Arithmetic Mean"</code> or <code>xlab = "Median"</code> when <code>plot = "boot"</code> .
ylab	a character string indicating the name argument in the <code>scale_y_continuous</code> function for labeling the y-axis. The default setting is <code>ylab = "Arithmetic Mean"</code> or <code>ylab = "Median"</code> when <code>plot = "ci"</code> and <code>ylab = "Probability Density, f(x)"</code> when <code>plot = "boot"</code> .
xlim	a numeric vector with two elements indicating the limits argument in the <code>scale_x_continuous</code> function for controlling the scale range of the x-axis.
ylim	a numeric vector with two elements indicating the limits argument in the <code>scale_y_continuous</code> function for controlling the scale range of the y-axis.
xbreaks	a numeric vector indicating the breaks argument in the <code>scale_x_continuous</code> function for controlling the x-axis breaks. The default setting is <code>xbreaks = NULL</code> when <code>plot = "ci"</code> and <code>xbreaks = seq(-1, 1, by = 0.25)</code> when <code>plot = "boot"</code> .
ybreaks	a numeric vector indicating the breaks argument in the <code>scale_y_continuous</code> function for controlling the y-axis breaks. The default setting is <code>ybreaks = seq(-1, 1, by = 0.25)</code> when <code>plot = "ci"</code> and <code>ybreaks = NULL</code> when <code>plot = "boot"</code> .
axis.title.size	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the axis title, i.e., <code>theme(axis.title = element_text(size = axis.text.size))</code> .

<code>axis.text.size</code>	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the axis text, i.e., <code>theme(axis.text = element_text(size = axis.text.size))</code> .
<code>strip.text.size</code>	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the strip text, i.e., <code>theme(strip.text = element_text(size = strip.text.size))</code> .
<code>title</code>	a character string indicating the title argument in the <code>labs</code> function for the subtitle of the plot.
<code>subtitle</code>	a character string indicating the subtitle argument in the <code>labs</code> function for the subtitle of the plot.
<code>group.col</code>	a character vector indicating the color argument in the <code>scale_color_manual</code> and <code>scale_fill_manual</code> functions when specifying a grouping variable using the argument group.
<code>plot.margin</code>	a numeric vector with four elements indicating the <code>plot.margin</code> argument in the <code>theme</code> function controlling the plot margins. The default setting is <code>c(5.5, 5.5, 5.5, 5.5)</code> , but switches to <code>c(5.5, 5.5, -2.5, 5.5)</code> when specifying a grouping variable using the argument group.
<code>legend.title</code>	a character string indicating the color argument in the <code>labs</code> function for specifying the legend title when specifying a grouping variable using the argument group.
<code>legend.position</code>	a character string indicating the <code>legend.position</code> in the <code>theme</code> argument for controlling the position of the legend function when specifying a grouping variable using the argument group. By default, the legend is placed at the bottom the plot.
<code>legend.box.margin</code>	a numeric vector with four elements indicating the <code>legend.box.margin</code> argument in the <code>theme</code> function for controlling the margins around the full legend area when specifying a grouping variable using the argument group.
<code>facet.ncol</code>	a numeric value indicating the <code>ncol</code> argument in the <code>facet_wrap</code> function for controlling the number of columns when specifying a split variable using the argument <code>split</code> .
<code>facet.nrow</code>	a numeric value indicating the <code>nrow</code> argument in the <code>facet_wrap</code> function for controlling the number of rows when specifying a split variable using the argument <code>split</code> .
<code>facet.scales</code>	a character string indicating the <code>scales</code> argument in the <code>facet_wrap</code> function for controlling the scales shared across facets, i.e., <code>"fixed"</code> , <code>"free_x"</code> , <code>"free_y"</code> , or <code>"free"</code> (default) when specifying a split variable using the argument <code>split</code> .
<code>filename</code>	a character string indicating the filename argument including the file extension in the <code>ggsave</code> function. Note that one of <code>".eps"</code> , <code>".ps"</code> , <code>".tex"</code> , <code>".pdf"</code> (default), <code>".jpeg"</code> , <code>".tiff"</code> , <code>".png"</code> , <code>".bmp"</code> , <code>".svg"</code> or <code>".wmf"</code> needs to be specified as file extension in the file argument. Note that plots can only be saved when <code>plot = "ci"</code> or <code>plot = "boot"</code> .

width	a numeric value indicating the width argument (default is the size of the current graphics device) in the ggsave function.
height	a numeric value indicating the height argument (default is the size of the current graphics device) in the ggsave function.
units	a character string indicating the units argument (default is in) in the ggsave function.
dpi	a numeric value indicating the dpi argument (default is 600) in the ggsave function.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	list with the input specified in data, group, and split
args	specification of function arguments
boot	data frame with bootstrap replicates of the arithmetic mean of median when bootstrapping was requested
plot	ggplot2 object for plotting the results and the data frame used for plotting
result	result table

Note

Bootstrap confidence intervals are computed using the R package `boot` by Angelo Canty and Brian Ripley (2024).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Baguley, T. S. (2012). *Serious stats: A guide to advanced statistics for the behavioral sciences*. Palgrave Macmillan.
- Canty, A., & Ripley, B. (2024). *boot: Bootstrap R (S-Plus) Functions*. R package version 1.3-31.
- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.

See Also

`test.z`, `test.t`, `ci.mean.diff`, `ci.cor`, `ci.prop`, `ci.var`, `ci.sd`, `descript`

Examples

```
#-----
# Confidence Interval (CI) for the Arithmetic Mean

# Example 1a: Two-Sided 95% CI
ci.mean(mtcars)

# Example 1b: Two-Sided 95% Difference-Adjusted CI
ci.mean(mtcars, adjust = TRUE)

# Example 1c: Two-Sided 95% CI with known population standard deviation
ci.mean(mtcars, mpg, sigma = 6)

# Alternative specification without using the '...' argument
ci.mean(mtcars$mpg, sigma = 6)

#-----
# Confidence Interval (CI) for the Median

# Example 2a: Two-Sided 95% CI
ci.median(mtcars)

# Example 2b: One-Sided 99% CI
ci.median(mtcars, alternative = "less", conf.level = 0.99)

## Not run:
#-----
# Bootstrap Confidence Interval (CI)

# Example 3a: Bias-corrected (BC) percentile bootstrap CI
ci.mean(mtcars, boot = "bc")

# Example 3b: Bias-corrected and accelerated (BCa) bootstrap CI,
# 5000 bootstrap replications, set seed of the pseudo-random number generator
ci.mean(mtcars, boot = "bca", R = 5000, seed = 123)

#-----
# Grouping and Split Variable

# Example 4a: Grouping variable
ci.mean(mtcars, mpg, cyl, disp, group = "vs")

# Alternative specification without using the '...' argument
ci.mean(mtcars[, c("mpg", "cyl", "disp")], group = mtcars$vs)

# Example 4b: Split variable
ci.mean(mtcars, mpg, cyl, disp, split = "am")
```

```
# Alternative specification
ci.mean(mtcars[, c("mpg", "cyl", "disp")], split = mtcars$am)

# Example 4c: Grouping and split variable
ci.mean(mtcars, mpg, cyl, disp, group = "vs", split = "am")

# Alternative specification
ci.mean(mtcars[, c("mpg", "cyl", "disp")], group = mtcars$vs, split = mtcars$am)

#-----
# Write Output

# Example 5a: Text file
ci.mean(mtcars, write = "CI_Mean_Text.txt")

# Example 5b: Excel file
ci.mean(mtcars, write = "CI_Mean_Excel.xlsx")

#-----
# Plot Confidence Intervals

# Example 6a: Two-Sided 95
ci.mean(mtcars, disp, hp, plot = "ci")

# Example 6b: Grouping variable
ci.mean(mtcars, disp, hp, group = "vs", plot = "ci")

# Example 6c: Split variable
ci.mean(mtcars, disp, hp, split = "am", plot = "ci")

# Example 6d: Save plot as PDF file
ci.mean(mtcars, disp, hp, plot = "ci", saveplot = "CI_Mean.pdf",
        width = 9, height = 6)

# Example 6e: Save plot as PNG file
ci.mean(mtcars, disp, hp, plot = "ci", saveplot = "CI_Mean.png",
        width = 9, height = 6)

#-----
# Example 7: Plot Bootstrap Samples

# Example 7a: Two-Sided 95
ci.mean(mtcars, disp, hp, boot = "bc", plot = "boot")

# Example 7b: Grouping variable
ci.mean(mtcars, disp, hp, group = "vs", boot = "bc", plot = "boot")

# Example 7c: Split variable
ci.mean(mtcars, disp, hp, split = "am", boot = "bc", plot = "boot")

# Example 7d: Save plot as PDF file
ci.mean(mtcars, disp, hp, boot = "bc", plot = "boot", saveplot = "CI_Mean_Boot.pdf",
        width = 12, height = 7)
```

```
# Example 7e: Save plot as PNG file
ci.mean(mtcars, disp, hp, boot = "bc", plot = "boot", saveplot = "CI_Mean_Boot.png",
        width = 12, height = 7)

## End(Not run)
```

ci.mean.diff

Confidence Interval for the Difference in Arithmetic Means

Description

This function computes a confidence interval for the difference in arithmetic means in a one-sample, two-sample and paired-sample design with known or unknown population standard deviation or population variance for one or more variables, optionally by a grouping and/or split variable.

Usage

```
ci.mean.diff(x, ...)

## Default S3 method:
ci.mean.diff(x, y, mu = 0, sigma = NULL, sigma2 = NULL,
             var.equal = FALSE, paired = FALSE,
             alternative = c("two.sided", "less", "greater"),
             conf.level = 0.95, group = NULL, split = NULL, sort.var = FALSE,
             digits = 2, as.na = NULL, write = NULL, append = TRUE,
             check = TRUE, output = TRUE, ...)

## S3 method for class 'formula'
ci.mean.diff(formula, data, sigma = NULL, sigma2 = NULL,
             var.equal = FALSE, alternative = c("two.sided", "less", "greater"),
             conf.level = 0.95, group = NULL, split = NULL, sort.var = FALSE,
             na.omit = FALSE, digits = 2, as.na = NULL, write = NULL,
             append = TRUE, check = TRUE, output = TRUE, ...)
```

Arguments

x	a numeric vector of data values.
...	further arguments to be passed to or from methods.
y	a numeric vector of data values.
mu	a numeric value indicating the population mean under the null hypothesis. Note that the argument mu is only used when y = NULL.
sigma	a numeric vector indicating the population standard deviation(s) when computing confidence intervals for the difference in arithmetic means with known standard deviation(s). In case of independent samples, equal standard deviations are assumed when specifying one value for the argument sigma; when specifying

	two values for the argument <code>sigma</code> , unequal standard deviations are assumed. Note that either argument <code>sigma</code> or argument <code>sigma2</code> is specified and it is only possible to specify one value (i.e., equal variance assumption) or two values (i.e., unequal variance assumption) for the argument <code>sigma</code> even though multiple variables are specified in <code>x</code> .
<code>sigma2</code>	a numeric vector indicating the population variance(s) when computing confidence intervals for the difference in arithmetic means with known variance(s). In case of independent samples, equal variances are assumed when specifying one value for the argument <code>sigma2</code> ; when specifying two values for the argument <code>sigma</code> , unequal variances are assumed. Note that either argument <code>sigma</code> or argument <code>sigma2</code> is specified and it is only possible to specify one value (i.e., equal variance assumption) or two values (i.e., unequal variance assumption) for the argument <code>sigma</code> even though multiple variables are specified in <code>x</code> .
<code>var.equal</code>	logical: if TRUE, the population variance in the independent samples are assumed to be equal.
<code>paired</code>	logical: if TRUE, confidence interval for the difference of arithmetic means in paired samples is computed.
<code>alternative</code>	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.
<code>group</code>	a numeric vector, character vector or factor as grouping variable. Note that a grouping variable can only be used when computing confidence intervals with unknown population standard deviation and population variance.
<code>split</code>	a numeric vector, character vector or factor as split variable. Note that a split variable can only be used when computing confidence intervals with unknown population
<code>sort.var</code>	logical: if TRUE, output table is sorted by variables when specifying <code>group</code> .
<code>digits</code>	an integer value indicating the number of decimal places to be used.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to <code>x</code> , but not to <code>group</code> or <code>split</code> .
<code>write</code>	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension ".txt" specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.
<code>formula</code>	a formula of the form <code>y ~ group</code> for one outcome variable or <code>cbind(y1, y2, y3) ~ group</code> for more than one outcome variable where <code>y</code> is a numeric variable giving the data values and <code>group</code> a numeric variable, character variable or factor with two values or factor levels giving the corresponding groups.
<code>data</code>	a matrix or data frame containing the variables in the formula <code>formula</code> .
<code>na.omit</code>	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion) when specifying more than one outcome variable.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	list with the input specified in <code>x</code> , <code>group</code> , and <code>split</code>
<code>args</code>	specification of function arguments
<code>result</code>	result table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.

See Also

[test.z](#), [test.t](#), [ci.mean](#), [ci.median](#), [ci.prop](#), [ci.var](#), [ci.sd](#), [descript](#)

Examples

```
#-----
# One-sample design

# Example 1a: Two-Sided 95% CI for 'mpg'
# population mean = 20
ci.mean.diff(mtcars$mpg, mu = 20)

# Example 1a: One-Sided 95% CI for 'mpg'
# population mean = 20
ci.mean.diff(mtcars$mpg, mu = 20, alternative = "greater")

#-----
# Two-sample design

# Example 2a: Two-Sided 95% CI for 'mpg' by 'vs'
# unknown population variances, unequal variance assumption
ci.mean.diff(mpg ~ vs, data = mtcars)

# Example 2b: Two-Sided 95% CI for 'mpg' by 'vs'
# unknown population variances, equal variance assumption
ci.mean.diff(mpg ~ vs, data = mtcars, var.equal = TRUE)

# Example 2c: Two-Sided 95% CI for 'mpg' by 'vs'
# known population standard deviations, equal standard deviation assumption
ci.mean.diff(mpg ~ vs, data = mtcars, sigma = 4)
```

```

# Example 2d: Two-Sided 95% CI for 'mpg' by 'vs'
# known population standard deviations, unequal standard deviation assumption
ci.mean.diff(mpg ~ vs, data = mtcars, sigma = c(4, 5))

# Example 2e: Two-Sided 95% CI for 'mpg', 'cyl', and 'disp' by 'vs'
# unknown population variances, unequal variance assumption
ci.mean.diff(cbind(mpg, cyl, disp) ~ vs, data = mtcars)

# Example 2f: Two-Sided 95% CI for 'mpg', 'cyl', and 'disp' by 'vs'
# unknown population variances, unequal variance assumption,
# analysis by am separately
ci.mean.diff(cbind(mpg, cyl, disp) ~ vs, data = mtcars, group = mtcars$am)

# Example 2g: Two-Sided 95% CI for 'mpg', 'cyl', and 'disp' by 'vs'
# unknown population variances, unequal variance assumption,
# split analysis by am
ci.mean.diff(cbind(mpg, cyl, disp) ~ vs, data = mtcars, split = mtcars$am)

# Example 2h: Two-Sided 95% CI for the mean difference between 'group1' and 'group2'
# unknown population variances, unequal variance assumption
group1 <- c(3, 1, 4, 2, 5, 3, 6, 7)
group2 <- c(5, 2, 4, 3, 1)

ci.mean.diff(group1, group2)

#-----
# Paired-sample design

dat.p <- data.frame(pre = c(1, 3, 2, 5, 7, 6), post = c(2, 2, 1, 6, 8, 9),
                    group = c(1, 1, 1, 2, 2, 2))

# Example 3a: Two-Sided 95% CI for the mean difference in 'pre' and 'post'
# unknown population variance of difference scores
ci.mean.diff(dat.p$pre, dat.p$post, paired = TRUE)

# Example 21: Two-Sided 95% CI for the mean difference in 'pre' and 'post'
# unknown population variance of difference scores
# analysis by group separately
ci.mean.diff(dat.p$pre, dat.p$post, paired = TRUE, group = dat.p$group)

# Example 22: Two-Sided 95% CI for the mean difference in 'pre' and 'post'
# unknown population variance of difference scores
# analysis by group separately
ci.mean.diff(dat.p$pre, dat.p$post, paired = TRUE, split = dat.p$group)

# Example 23: Two-Sided 95% CI for the mean difference in 'pre' and 'post'
# known population standard deviation of difference scores
ci.mean.diff(dat.p$pre, dat.p$post, sigma = 2, paired = TRUE)

```

Description

This function computes difference-adjusted Cousineau-Morey within-subject confidence interval for the arithmetic mean.

Usage

```
ci.mean.w(data, ..., adjust = TRUE,
           alternative = c("two.sided", "less", "greater"), conf.level = 0.95,
           na.omit = TRUE, digits = 2, as.na = NULL, write = NULL, append = TRUE,
           check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a data frame with numeric variables representing the levels of the within-subject factor, i.e., data are specified in wide-format (i.e., multivariate person level format).
<code>...</code>	an expression indicating the variable names in <code>data</code> , e.g., <code>ci.mean.w(dat, time1, time2, time3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>adjust</code>	logical: if TRUE (default), difference-adjustment for the Cousineau-Morey within-subject confidence intervals is applied.
<code>alternative</code>	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.
<code>na.omit</code>	logical: if TRUE (default), incomplete cases are removed before conducting the analysis (i.e., listwise deletion).
<code>digits</code>	an integer value indicating the number of decimal places to be used.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>write</code>	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension ".txt" specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

The Cousineau within-subject confidence interval (CI, Cousineau, 2005) is an alternative to the Loftus-Masson within-subject CI (Loftus & Masson, 1994) that does not assume sphericity or homogeneity of covariances. This approach removes individual differences by normalizing the raw scores using participant-mean centering and adding the grand mean back to every score:

$$Y'_{ij} = Y_{ij} - \hat{\mu}_i + \hat{\mu}_{grand}$$

where Y'_{ij} is the score of the i th participant in condition j (for $i = 1$ to n), $\hat{\mu}_i$ is the mean of participant i across all J levels (for $j = 1$ to J), and $\hat{\mu}_{grand}$ is the grand mean.

Morey (2008) pointed out that Cousineau's (2005) approach produces intervals that are consistently too narrow due to inducing a positive covariance between normalized scores within a condition introducing bias into the estimate of the sample variances. The degree of bias is proportional to the number of means and can be removed by rescaling the confidence interval by a factor of $\sqrt{J-1}/J$:

$$\hat{\mu}_j \pm t_{n-1, 1-\alpha/2} \sqrt{\frac{J}{J-1}} \hat{\sigma}'_{\hat{\mu}_j}$$

where $\hat{\sigma}'_{\hat{\mu}_j}$ is the standard error of the mean computed from the normalized scores of the j th factor level.

Baguley (2012) pointed out that the Cousineau-Morey interval is larger than that for a difference in means by a factor of $\sqrt{2}$ leading to a misinterpretation of these intervals that overlap of 95% confidence intervals around individual means indicates that a 95% confidence interval for the difference in means would include zero. Hence, following adjustment to the Cousineau-Morey interval was proposed:

$$\hat{\mu}_j \pm \frac{\sqrt{2}}{2} (t_{n-1, 1-\alpha/2} \sqrt{\frac{J}{J-1}} \hat{\sigma}'_{\hat{\mu}_j})$$

The adjusted Cousineau-Morey interval is informative about the pattern of differences between means and is computed by default (i.e., `adjust = TRUE`).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame used for the current analysis
<code>args</code>	specification of function arguments
<code>result</code>	result table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Baguley, T. (2012). Calculating and graphing within-subject confidence intervals for ANOVA. *Behavior Research Methods*, 44, 158-175. <https://doi.org/10.3758/s13428-011-0123-7>

Cousineau, D. (2005) Confidence intervals in within-subject designs: A simpler solution to Loftus and Masson's Method. *Tutorials in Quantitative Methods for Psychology*, 1, 42-45. <https://doi.org/10.20982/tqmp.01.1.p042>

Loftus, G. R., and Masson, M. E. J. (1994). Using confidence intervals in within-subject designs. *Psychonomic Bulletin and Review*, 1, 476-90. <https://doi.org/10.3758/BF03210951>

Morey, R. D. (2008). Confidence intervals from normalized data: A correction to Cousineau. *Tutorials in Quantitative Methods for Psychology*, 4, 61-4. <https://doi.org/10.20982/tqmp.01.1.p042>

See Also

[aov.w](#), [test.z](#), [test.t](#), [ci.mean.diff](#), [ci.median](#), [ci.prop](#), [ci.var](#), [ci.sd](#), [descript](#)

Examples

```
dat <- data.frame(time1 = c(3, 2, 1, 4, 5, 2, 3, 5, 6, 7),
                  time2 = c(4, 3, 6, 5, 8, 6, 7, 3, 4, 5),
                  time3 = c(1, 2, 2, 3, 6, 5, 1, 2, 4, 6))

# Example 1: Difference-adjusted Cousineau-Morey confidence intervals
ci.mean.w(dat)

# Example 2: Cousineau-Morey confidence intervals
ci.mean.w(dat, adjust = FALSE)

## Not run:
# Example 3: Write results into a text file
ci.mean.w(dat, write = "WS_Confidence_Interval.txt")
## End(Not run)
```

ci.prop

(Bootstrap) Confidence Intervals for Proportions

Description

This function computes and plots confidence intervals for proportions, optionally by a grouping and/or split variable. The function also supports three types of bootstrap confidence intervals (e.g., bias-corrected (BC) percentile bootstrap or bias-corrected and accelerated (BCa) bootstrap confidence intervals) and plots the bootstrap samples with histograms and density curves.

Usage

```
ci.prop(data, ..., method = c("wald", "wilson"),
        boot = c("none", "perc", "bc", "bca"), R = 1000, seed = NULL,
        alternative = c("two.sided", "less", "greater"), conf.level = 0.95,
        group = NULL, split = NULL, sort.var = FALSE, na.omit = FALSE, digits = 3,
        as.na = NULL, plot = c("none", "ci", "boot"), point.size = 2.5,
        point.shape = 19, errorbar.width = 0.3, dodge.width = 0.5, hist = TRUE,
        binwidth = NULL, bins = NULL, hist.alpha = 0.4, fill = "gray85",
        density = TRUE, density.col = "#0072B2", density.linewidth = 0.5,
        density.linetype = "solid", point = TRUE, point.col = "#CC79A7",
        point.linewidth = 0.6, point.linetype = "solid", ci = TRUE, ci.col = "black",
        ci.linewidth = 0.6, ci.linetype = "dashed", line = FALSE, intercept = 0.5,
        linetype = "solid", line.col = "gray65", xlab = NULL, ylab = NULL,
        xlim = NULL, ylim = NULL, xbreaks = ggplot2::waiver(),
        ybreaks = ggplot2::waiver(), axis.title.size = 11, axis.text.size = 10,
        strip.text.size = 11, title = NULL, subtitle = NULL, group.col = NULL,
        plot.margin = NA, legend.title = "",
```

```

legend.position = c("right", "top", "left", "bottom", "none"),
legend.box.margin = c(-10, 0, 0, 0), facet.ncol = NULL, facet.nrow = NULL,
facet.scales = "free_y", filename = NULL, width = NA, height = NA,
units = c("in", "cm", "mm", "px"), dpi = 600, write = NULL, append = TRUE,
check = TRUE, output = TRUE)

```

Arguments

<code>data</code>	a numeric vector or data frame with numeric variables with 0 and 1 values.
<code>...</code>	an expression indicating the variable names in data, e.g., <code>ci.prop(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>method</code>	a character string specifying the method for computing the confidence interval, must be one of "wald", or "wilson" (default).
<code>boot</code>	a character string specifying the type of bootstrap confidence intervals (CI), i.e., "none" (default) for not conducting bootstrapping, "perc", for the percentile bootstrap CI "bc" (default) for the bias-corrected (BC) percentile bootstrap CI (without acceleration), and "bca" for the bias-corrected and accelerated (BCa) bootstrap CI, see 'Details' in the ci.cor function.
<code>R</code>	a numeric value indicating the number of bootstrap replicates (default is 1000).
<code>seed</code>	a numeric value specifying seeds of the pseudo-random numbers used in the bootstrap algorithm when conducting bootstrapping.
<code>alternative</code>	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.
<code>group</code>	either a character string indicating the variable name of the grouping variable in data, or a vector representing the grouping variable.
<code>split</code>	either a character string indicating the variable name of the split variable in data, or a vector representing the split variable.
<code>sort.var</code>	logical: if TRUE, output table is sorted by variables when specifying group.
<code>na.omit</code>	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion) when specifying more than one outcome variable.
<code>digits</code>	an integer value indicating the number of decimal places to be used.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to data, but not to group or split.
<code>plot</code>	a character string indicating the type of the plot to display, i.e., "none" (default) for not displaying any plots, "ci" for displaying confidence intervals for the proportion, "boot" for displaying bootstrap samples with histograms and density curves when the argument "boot" is other than "none".
<code>point.size</code>	a numeric value indicating the size argument in the <code>geom_point</code> function for controlling the size of points when plotting confidence intervals (plot = "ci").
<code>point.shape</code>	a numeric value between 0 and 25 or a character string as plotting symbol indicating the shape argument in the <code>geom_point</code> function for controlling the symbols of points when plotting confidence intervals (plot = "ci").

errorbar.width	a numeric value indicating the width argument in the <code>geom_errorbar</code> function for controlling the width of the whiskers in the <code>geom_errorbar</code> function when plotting confidence intervals (plot = "ci").
dodge.width	a numeric value indicating the width argument controlling the width of the geom elements to be dodged when specifying a grouping variable using the argument group when plotting confidence intervals (plot = "ci").
hist	logical: if TRUE (default), histograms are drawn when plotting bootstrap samples (plot = "boot").
binwidth	a numeric value or a function for specifying the binwidth argument in the <code>geom_histogram</code> function for controlling the width of the bins when plotting bootstrap samples (plot = "boot").
bins	a numeric value for specifying the bins argument in the <code>geom_histogram</code> function for controlling the number of bins when plotting bootstrap samples (plot = "boot").
hist.alpha	a numeric value between 0 and 1 for specifying the alpha argument in the <code>geom_histogram</code> function for controlling the opacity of the bars when plotting bootstrap samples (plot = "boot").
fill	a character string specifying the fill argument in the <code>geom_histogram</code> function controlling the fill aesthetic when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
density	logical: if TRUE (default), density curves are drawn when plotting bootstrap samples (plot = "boot").
density.col	a character string specifying the color argument in the <code>geom_density</code> function controlling the color of the density curves when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
density.linewidth	a numeric value specifying the linewidth argument in the <code>geom_density</code> function controlling the line width of the density curves when plotting bootstrap samples (plot = "boot").
density.linetype	a numeric value or character string specifying the linetype argument in the <code>geom_density</code> function controlling the line type of the density curves when plotting bootstrap samples (plot = "boot").
point	logical: if TRUE (default), vertical lines representing the point estimate of the proportion are drawn when plotting bootstrap samples (plot = "boot").
point.col	a character string specifying the color argument in the <code>geom_vline</code> function for controlling the color of the vertical line displaying the proportion when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
point.linewidth	a numeric value specifying the linewidth argument in the <code>geom_vline</code> function for controlling the line width of the vertical line displaying proportions when plotting bootstrap samples (plot = "boot").

point.linetype	a numeric value or character string specifying the linetype argument in the geom_vline function controlling the line type of the vertical line displaying proportions when plotting bootstrap samples (plot = "boot").
ci	logical: if TRUE (default), vertical lines representing the bootstrap confidence intervals of proportions are drawn when plotting bootstrap samples (plot = "boot").
ci.col	character string specifying the color argument in the geom_vline function for controlling the color of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
ci.linewidth	a numeric value specifying the linewidth argument in the geom_vline function for controlling the line width of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (plot = "boot").
ci.linetype	a numeric value or character string specifying the linetype argument in the geom_vline function controlling the line type of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (plot = "boot").
line	logical: if TRUE, a horizontal line is drawn when plot = "ci" or a vertical line is drawn when plot = "boot"
intercept	a numeric value indicating the yintercept or xintercept argument in the geom_hline or geom_vline function controlling the position of the horizontal or vertical line when plot = "ci" and line = TRUE or when plot = "boot" and line = TRUE. By default, the horizontal or vertical line is drawn at 0.
linetype	a character string indicating the linetype argument in the geom_hline or geom_vline function controlling the line type of the horizontal or vertical line (default is linetype = "dashed").
line.col	a character string indicating the color argument in the geom_hline or geom_vline function for controlling the color of the horizontal or vertical line.
xlab	a character string indicating the name argument in the scale_x_continuous function for labeling the x-axis. The default setting is xlab = NULL when plot = "ci" and xlab = "Proportion" when plot = "boot".
ylab	a character string indicating the name argument in the scale_y_continuous function for labeling the y-axis. The default setting is ylab = "Proportion" when plot = "ci" and ylab = "Probability Density, f(x)" when plot = "boot".
xlim	a numeric vector with two elements indicating the limits argument in the scale_x_continuous function for controlling the scale range of the x-axis. The default setting is xlim = NULL when plot = "ci" and xlim = c(0, 1) when plot = "boot".
ylim	a numeric vector with two elements indicating the limits argument in the scale_y_continuous function for controlling the scale range of the y-axis. The default setting is ylim = c(0, 1) when plot = "ci" and ylim = NULL when plot = "boot".
xbreaks	a numeric vector indicating the breaks argument in the scale_x_continuous function for controlling the x-axis breaks. The default setting is xbreaks = NULL when plot = "ci" and xbreaks = seq(-1, 1, by = 0.25) when plot = "boot".

ybreaks	a numeric vector indicating the breaks argument in the <code>scale_y_continuous</code> function for controlling the y-axis breaks. The default setting is <code>ybreaks = seq(-1, 1, by = 0.25)</code> when <code>plot = "ci"</code> and <code>ybreaks = NULL</code> when <code>plot = "boot"</code> .
axis.title.size	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the axis title, i.e., <code>theme(axis.title = element_text(size = axis.text.size))</code> .
axis.text.size	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the axis text, i.e., <code>theme(axis.text = element_text(size = axis.text.size))</code> .
strip.text.size	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the strip text, i.e., <code>theme(strip.text = element_text(size = strip.text.size))</code> .
title	a character string indicating the title argument in the <code>labs</code> function for the subtitle of the plot.
subtitle	a character string indicating the subtitle argument in the <code>labs</code> function for the subtitle of the plot.
group.col	a character vector indicating the color argument in the <code>scale_color_manual</code> and <code>scale_fill_manual</code> functions when specifying a grouping variable using the argument <code>group</code> .
plot.margin	a numeric vector with four elements indicating the <code>plot.margin</code> argument in the <code>theme</code> function controlling the plot margins. The default setting is <code>c(5.5, 5.5, 5.5, 5.5)</code> , but switches to <code>c(5.5, 5.5, -2.5, 5.5)</code> when specifying a grouping variable using the argument <code>group</code> .
legend.title	a character string indicating the color argument in the <code>labs</code> function for specifying the legend title when specifying a grouping variable using the argument <code>group</code> .
legend.position	a character string indicating the <code>legend.position</code> in the <code>theme</code> argument for controlling the position of the legend function when specifying a grouping variable using the argument <code>group</code> . By default, the legend is placed at the bottom the plot.
legend.box.margin	a numeric vector with four elements indicating the <code>legend.box.margin</code> argument in the <code>theme</code> function for controlling the margins around the full legend area when specifying a grouping variable using the argument <code>group</code> .
facet.ncol	a numeric value indicating the <code>ncol</code> argument in the <code>facet_wrap</code> function for controlling the number of columns when specifying a split variable using the argument <code>split</code> .
facet.nrow	a numeric value indicating the <code>nrow</code> argument in the <code>facet_wrap</code> function for controlling the number of rows when specifying a split variable using the argument <code>split</code> .

facet.scales	a character string indicating the scales argument in the facet_wrap function for controlling the scales shared across facets, i.e., "fixed", "free_x", "free_y", or "free" (default) when specifying a split variable using the argument split.
filename	a character string indicating the filename argument including the file extension in the ggsave function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the file argument. Note that plots can only be saved when plot = "ci" or plot = "boot".
width	a numeric value indicating the width argument (default is the size of the current graphics device) in the ggsave function.
height	a numeric value indicating the height argument (default is the size of the current graphics device) in the ggsave function.
units	a character string indicating the units argument (default is in) in the ggsave function.
dpi	a numeric value indicating the dpi argument (default is 600) in the ggsave function.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

The Wald confidence interval which is based on the normal approximation to the binomial distribution are computed by specifying method = "wald", while the Wilson (1927) confidence interval (aka Wilson score interval) is requested by specifying method = "wilson". By default, Wilson confidence interval is computed which have been shown to be reliable in small samples of $n = 40$ or less, and larger samples of $n > 40$ (Brown, Cai & DasGupta, 2001), while the Wald confidence intervals is inadequate in small samples and when p is near 0 or 1 (Agresti & Coull, 1998).

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	list with the input specified in ..., data, group, and split
args	specification of function arguments
boot	data frame with bootstrap replicates of the aproportion when bootstrapping was requested
plot	ggplot2 object for plotting the results and the data frame used for plotting
result	result table

Note

Bootstrap confidence intervals are computed using the R package `boot` by Angelo Canty and Brian Ripley (2024).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Agresti, A. & Coull, B.A. (1998). Approximate is better than "exact" for interval estimation of binomial proportions. *American Statistician*, 52, 119-126.
- Brown, L. D., Cai, T. T., & DasGupta, A., (2001). Interval estimation for a binomial proportion. *Statistical Science*, 16, 101-133.
- Canty, A., & Ripley, B. (2024). *boot: Bootstrap R (S-Plus) Functions*. R package version 1.3-31.
- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.
- Wilson, E. B. (1927). Probable inference, the law of succession, and statistical inference. *Journal of the American Statistical Association*, 22, 209-212.

See Also

[ci.prop](#), [ci.prop.diff](#), [ci.median](#), [ci.prop.diff](#), [ci.cor](#), [ci.var](#), [ci.sd](#), [descript](#)

Examples

```
#-----
# Confidence Interval (CI) for proportions

# Example 1a: Two-Sided 95% CI
ci.prop(mtcars, vs, am)

# Alternative specification without using the '...' argument
ci.prop(mtcars[, c("vs", "am")])

# Example 1b: One-Sided 95% CI using Wald method
ci.prop(mtcars, vs, am, method = "wald", alternative = "less")

## Not run:
#-----
# Bootstrap Confidence Interval (CI)

# Example 2a: Bias-corrected (BC) percentile bootstrap CI
ci.prop(mtcars, vs, am, boot = "bc")

# Example 2b: Bias-corrected and accelerated (BCa) bootstrap CI,
# 5000 bootstrap replications, set seed of the pseudo-random number generator
ci.prop(mtcars, vs, am, boot = "bca", R = 5000, seed = 123)
```

```
#-----
# Grouping and Split Variable

# Example 3a: Grouping variable
ci.prop(mtcars, vs, group = "am")

# Alternative specification without using the '...' argument
ci.prop(mtcars$vs, group = mtcars$am)

# Example 3b: Split variable
ci.prop(mtcars, vs, split = "am")

# Alternative specification without using the '...' argument
ci.prop(mtcars$vs, split = mtcars$am)

# Example 3c: Grouping and split variable
ci.prop(mtcars, vs, group = "am", split = "cyl")

# Alternative specification without using the '...' argument
ci.prop(mtcars$vs, group = mtcars$am, split = mtcars$cyl)

#-----
# Write Output

# Example 4a: Text file
ci.prop(mtcars, vs, am, write = "CI_Prop_Text.txt")

# Example 4b: Excel file
ci.prop(mtcars, vs, am, write = "CI_Prop_Excel.xlsx")

#-----
# Plot Confidence Intervals

# Example 5a: Two-Sided 95
ci.prop(mtcars, vs, am, plot = "ci")

# Example 5b: Grouping variable
ci.prop(mtcars, vs, am, group = "am", plot = "ci")

# Example 5c: Split variable
ci.prop(mtcars, vs, am, split = "am", plot = "ci")

# Example 5d: Save plot as PDF file
ci.prop(mtcars, vs, am, plot = "ci", saveplot = "CI_Prop.pdf",
        width = 9, height = 6)

# Example 5e: Save plot as PNG file
ci.prop(mtcars, vs, am, plot = "ci", saveplot = "CI_Prop.png",
        width = 9, height = 6)

#-----
# Plot Bootstrap Samples
```

```
# Example 6a: Two-Sided 95
ci.prop(mtcars, vs, am, boot = "bc", plot = "boot")

# Example 6b: Grouping variable
ci.prop(mtcars, vs, am, group = "am", boot = "bc", plot = "boot")

# Example 6c: Split variable
ci.prop(mtcars, vs, am, split = "am", boot = "bc", plot = "boot")

# Example 6d: Save plot as PDF file
ci.prop(mtcars, vs, am, boot = "bc", plot = "boot",
        saveplot = "CI_Prop_Boot.pdf", width = 9, height = 6)

# Example 6e: Save plot as PNG file
ci.prop(mtcars, vs, am, boot = "bc", plot = "boot",
        saveplot = "CI_Prop_Boot.png", width = 9, height = 6)

## End(Not run)
```

ci.prop.diff

Confidence Interval for the Difference in Proportions

Description

This function computes a confidence interval for the difference in proportions in a two-sample and paired-sample design for one or more variables, optionally by a grouping and/or split variable.

Usage

```
ci.prop.diff(x, ...)

## Default S3 method:
ci.prop.diff(x, y, method = c("wald", "newcombe"), paired = FALSE,
             alternative = c("two.sided", "less", "greater"), conf.level = 0.95,
             group = NULL, split = NULL, sort.var = FALSE, digits = 2,
             as.na = NULL, write = NULL, append = TRUE,
             check = TRUE, output = TRUE, ...)

## S3 method for class 'formula'
ci.prop.diff(formula, data, method = c("wald", "newcombe"),
             alternative = c("two.sided", "less", "greater"), conf.level = 0.95,
             group = NULL, split = NULL, sort.var = FALSE, na.omit = FALSE,
             digits = 2, as.na = NULL, write = NULL, append = TRUE,
             check = TRUE, output = TRUE, ...)
```

Arguments

x a numeric vector with 0 and 1 values.

...	further arguments to be passed to or from methods.
y	a numeric vector with 0 and 1 values.
method	a character string specifying the method for computing the confidence interval, must be one of "wald", or "newcombe" (default).
paired	logical: if TRUE, confidence interval for the difference of proportions in paired samples is computed.
alternative	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
group	a numeric vector, character vector or factor as grouping variable. Note that a grouping variable can only be used when computing confidence intervals with unknown population standard deviation and population variance.
split	a numeric vector, character vector or factor as split variable. Note that a split variable can only be used when computing confidence intervals with unknown population standard deviation and population variance.
sort.var	logical: if TRUE, output table is sorted by variables when specifying group.
digits	an integer value indicating the number of decimal places to be used.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that as.na() function is only applied to x, but not to group or split.
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.
formula	a formula of the form $y \sim \text{group}$ for one outcome variable or $\text{cbind}(y1, y2, y3) \sim \text{group}$ for more than one outcome variable where y is a numeric variable with 0 and 1 values and group a numeric variable, character variable or factor with two values or factor levels giving the corresponding group.
data	a matrix or data frame containing the variables in the formula formula.
na.omit	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion) when specifying more than one outcome variable.

Details

The Wald confidence interval which is based on the normal approximation to the binomial distribution are computed by specifying `method = "wald"`, while the Newcombe Hybrid Score interval (Newcombe, 1998a; Newcombe, 1998b) is requested by specifying `method = "newcombe"`. By default, Newcombe Hybrid Score interval is computed which have been shown to be reliable in small samples (less than $n = 30$ in each sample) as well as moderate to larger samples ($n > 30$ in each sample) and with proportions close to 0 or 1, while the Wald confidence intervals does not perform well unless the sample size is large (Fagerland, Lydersen & Laake, 2011).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	list with the input specified in <code>x</code> , <code>group</code> , and <code>split</code>
<code>args</code>	specification of function arguments
<code>result</code>	result table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Fagerland, M. W., Lydersen S., & Laake, P. (2011) Recommended confidence intervals for two independent binomial proportions. *Statistical Methods in Medical Research*, 24, 224-254.
- Newcombe, R. G. (1998a). Interval estimation for the difference between independent proportions: Comparison of eleven methods. *Statistics in Medicine*, 17, 873-890.
- Newcombe, R. G. (1998b). Improved confidence intervals for the difference between binomial proportions based on paired data. *Statistics in Medicine*, 17, 2635-2650.
- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.

See Also

[ci.prop](#), [ci.mean](#), [ci.mean.diff](#), [ci.median](#), [ci.var](#), [ci.sd](#), [descript](#)

Examples

```
#-----
# Two-sample design

# Example 1a: Two-Sided 95% CI for 'vs' by 'am'
# Newcombes Hybrid Score interval
ci.prop.diff(vs ~ am, data = mtcars)

# Example 1b: Two-Sided 95% CI for 'vs' by 'am'
# Wald CI
ci.prop.diff(vs ~ am, data = mtcars, method = "wald")

# Example 1c: Two-Sided 95% CI for the difference in proportions
# Newcombes Hybrid Score interval
ci.prop.diff(c(0, 1, 1, 0, 0, 1, 0, 1), c(1, 1, 1, 0, 0))

#-----
# Paired-sample design
```

```

dat.p <- data.frame(pre = c(0, 1, 1, 0, 1), post = c(1, 1, 0, 1, 1))

# Example 2a: Two-Sided 95% CI for the difference in proportions 'pre' and 'post'
# Newcombes Hybrid Score interval
ci.prop.diff(dat.p$pre, dat.p$post, paired = TRUE)

# Example 2b: Two-Sided 95% CI for the difference in proportions 'pre' and 'post'
# Wald CI
ci.prop.diff(dat.p$pre, dat.p$post, method = "wald", paired = TRUE)

```

ci.var	<i>(Bootstrap) Confidence Intervals for Variances and Standard Deviations</i>
--------	---

Description

The function `ci.var` computes and plots confidence intervals for variances, and the function `ci.sd` computes confidence intervals for the standard deviations, optionally by a grouping and/or split variable. These functions also supports three types of bootstrap confidence intervals (e.g., bias-corrected (BC) percentile bootstrap or bias-corrected and accelerated (BCa) bootstrap confidence intervals) and plots the bootstrap samples with histograms and density curves.

Usage

```

ci.var(data, ..., method = c("chisq", "bonett"),
       boot = c("none", "perc", "bc", "bca"), R = 1000, seed = NULL,
       alternative = c("two.sided", "less", "greater"),
       conf.level = 0.95, group = NULL, split = NULL, sort.var = FALSE,
       na.omit = FALSE, digits = 2, as.na = NULL,
       plot = c("none", "ci", "boot"), point.size = 2.5, point.shape = 19,
       errorbar.width = 0.3, dodge.width = 0.5, hist = TRUE,
       binwidth = NULL, bins = NULL, hist.alpha = 0.4, fill = "gray85", density = TRUE,
       density.col = "#0072B2", density.linewidth = 0.5, density.linetype = "solid",
       point = TRUE, point.col = "#CC79A7", point.linewidth = 0.6,
       point.linetype = "solid", ci = TRUE, ci.col = "black",
       ci.linewidth = 0.6, ci.linetype = "dashed", line = FALSE, intercept = 0,
       linetype = "solid", line.col = "gray65", xlab = NULL, ylab = NULL,
       xlim = NULL, ylim = NULL, xbreaks = ggplot2::waiver(), ybreaks = ggplot2::waiver(),
       axis.title.size = 11, axis.text.size = 10, strip.text.size = 11, title = NULL,
       subtitle = NULL, group.col = NULL, plot.margin = NA, legend.title = "",
       legend.position = c("right", "top", "left", "bottom", "none"),
       legend.box.margin = c(-10, 0, 0, 0), facet.ncol = NULL, facet.nrow = NULL,
       facet.scales = "free", filename = NULL, width = NA, height = NA,
       units = c("in", "cm", "mm", "px"), dpi = 600, write = NULL, append = TRUE,
       check = TRUE, output = TRUE)

ci.sd(data, ..., method = c("chisq", "bonett"),
      boot = c("none", "perc", "bc", "bca"), R = 1000, seed = NULL,

```

```

alternative = c("two.sided", "less", "greater"),
conf.level = 0.95, group = NULL, split = NULL, sort.var = FALSE,
na.omit = FALSE, digits = 2, as.na = NULL,
plot = c("none", "ci", "boot"), point.size = 2.5, point.shape = 19,
errorbar.width = 0.3, dodge.width = 0.5, hist = TRUE,
binwidth = NULL, bins = NULL, hist.alpha = 0.4, fill = "gray85", density = TRUE,
density.col = "#0072B2", density.linewidth = 0.5, density.linetype = "solid",
point = TRUE, point.col = "#CC79A7", point.linewidth = 0.6,
point.linetype = "solid", ci = TRUE, ci.col = "black",
ci.linewidth = 0.6, ci.linetype = "dashed", line = FALSE, intercept = 0,
linetype = "solid", line.col = "gray65", xlab = NULL, ylab = NULL,
xlim = NULL, ylim = NULL, xbreaks = ggplot2::waiver(), ybreaks = ggplot2::waiver(),
axis.title.size = 11, axis.text.size = 10, strip.text.size = 11, title = NULL,
subtitle = NULL, group.col = NULL, plot.margin = NA, legend.title = "",
legend.position = c("right", "top", "left", "bottom", "none"),
legend.box.margin = c(-10, 0, 0, 0), facet.ncol = NULL, facet.nrow = NULL,
facet.scales = "free", filename = NULL, width = NA, height = NA,
units = c("in", "cm", "mm", "px"), dpi = 600, write = NULL,
append = TRUE, check = TRUE, output = TRUE)

```

Arguments

<code>data</code>	a numeric vector or data frame with numeric variables, i.e., factors and character variables are excluded from data before conducting the analysis.
<code>...</code>	an expression indicating the variable names in data, e.g., <code>ci.var(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>method</code>	a character string specifying the method for computing the confidence interval, must be one of <code>"chisq"</code> , or <code>"bonett"</code> (default).
<code>boot</code>	a character string specifying the type of bootstrap confidence intervals (CI), i.e., <code>"none"</code> (default) for not conducting bootstrapping, <code>"perc"</code> , for the percentile bootstrap CI <code>"bc"</code> (default) for the bias-corrected (BC) percentile bootstrap CI (without acceleration), and <code>"bca"</code> for the bias-corrected and accelerated (BCa) bootstrap CI, see 'Details' in the ci.cor function.
<code>R</code>	a numeric value indicating the number of bootstrap replicates (default is 1000).
<code>seed</code>	a numeric value specifying seeds of the pseudo-random numbers used in the bootstrap algorithm when conducting bootstrapping.
<code>alternative</code>	a character string specifying the alternative hypothesis, must be one of <code>"two.sided"</code> (default), <code>"greater"</code> or <code>"less"</code> .
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.
<code>group</code>	either a character string indicating the variable name of the grouping variable in data, or a vector representing the grouping variable.
<code>split</code>	either a character string indicating the variable name of the split variable in 'data', or a vector representing the split variable.
<code>sort.var</code>	logical: if TRUE, output table is sorted by variables when specifying group.

na.omit	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion) when specifying more than one outcome variable.
digits	an integer value indicating the number of decimal places to be used.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that as.na() function is only applied to data, but not to group or split.
plot	a character string indicating the type of the plot to display, i.e., "none" (default) for not displaying any plots, "ci" for displaying confidence intervals for variances or standard deviations, "boot" for displaying bootstrap samples with histograms and density curves when the argument "boot" is other than "none".
point.size	a numeric value indicating the size argument in the geom_point function for controlling the size of points when plotting confidence intervals (plot = "ci").
point.shape	a numeric value between 0 and 25 or a character string as plotting symbol indicating the shape argument in the geom_point function for controlling the symbols of points when plotting confidence intervals (plot = "ci").
errorbar.width	a numeric value indicating the width argument in the geom_errorbar function for controlling the width of the whiskers in the geom_errorbar function when plotting confidence intervals (plot = "ci").
dodge.width	a numeric value indicating the width argument controlling the width of the geom elements to be dodged when specifying a grouping variable using the argument group and plotting confidence intervals (plot = "ci").
hist	logical: if TRUE (default), histograms are drawn when plotting bootstrap samples (plot = "boot").
binwidth	a numeric value or a function for specifying the binwidth argument in the geom_histogram function for controlling the width of the bins when plotting bootstrap samples (plot = "boot").
bins	a numeric value for specifying the bins argument in the geom_histogram function for controlling the number of bins when plotting bootstrap samples (plot = "boot").
hist.alpha	a numeric value between 0 and 1 for specifying the alpha argument in the geom_histogram function for controlling the opacity of the bars when plotting bootstrap samples (plot = "boot").
fill	a character string specifying the fill argument in the geom_histogram function controlling the fill aesthetic when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
density	logical: if TRUE (default), density curves are drawn when plotting bootstrap samples (plot = "boot").
density.col	a character string specifying the color argument in the geom_density function controlling the color of the density curves when plotting bootstrap samples (plot = "boot"). Note that this argument applied only when no grouping variable was specified group = NULL.
density.linewidth	a numeric value specifying the linewidth argument in the geom_density function controlling the line width of the density curves when plotting bootstrap samples (plot = "boot").

<code>density.linetype</code>	a numeric value or character string specifying the <code>linetype</code> argument in the <code>geom_density</code> function controlling the line type of the density curves when plotting bootstrap samples (<code>plot = "boot"</code>).
<code>point</code>	logical: if TRUE (default), vertical lines representing the point estimate of the variance or standard deviation are drawn when plotting bootstrap samples (<code>plot = "boot"</code>).
<code>point.col</code>	a character string specifying the <code>color</code> argument in the <code>geom_vline</code> function for controlling the color of the vertical line displaying the variance or standard deviation when plotting bootstrap samples (<code>plot = "boot"</code>). Note that this argument applied only when no grouping variable was specified <code>group = NULL</code> .
<code>point.linewidth</code>	a numeric value specifying the <code>linewidth</code> argument in the <code>geom_vline</code> function for controlling the line width of the vertical line displaying the variance or standard deviation when plotting bootstrap samples (<code>plot = "boot"</code>).
<code>point.linetype</code>	a numeric value or character string specifying the <code>linetype</code> argument in the <code>geom_vline</code> function controlling the line type of the vertical line displaying the variance or standard deviation when plotting bootstrap samples (<code>plot = "boot"</code>).
<code>ci</code>	logical: if TRUE (default), vertical lines representing the bootstrap confidence intervals of the variance or standard deviation are drawn when plotting bootstrap samples (<code>plot = "boot"</code>).
<code>ci.col</code>	character string specifying the <code>color</code> argument in the <code>geom_vline</code> function for controlling the color of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (<code>plot = "boot"</code>). Note that this argument applied only when no grouping variable was specified <code>group = NULL</code> .
<code>ci.linewidth</code>	a numeric value specifying the <code>linewidth</code> argument in the <code>geom_vline</code> function for controlling the line width of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (<code>plot = "boot"</code>).
<code>ci.linetype</code>	a numeric value or character string specifying the <code>linetype</code> argument in the <code>geom_vline</code> function controlling the line type of the vertical line displaying bootstrap confidence intervals when plotting bootstrap samples (<code>plot = "boot"</code>).
<code>line</code>	logical: if TRUE, a horizontal line is drawn when <code>plot = "ci"</code> or a vertical line is drawn when <code>plot = "boot"</code>
<code>intercept</code>	a numeric value indicating the <code>yintercept</code> or <code>xintercept</code> argument in the <code>geom_hline</code> or <code>geom_vline</code> function controlling the position of the horizontal or vertical line when <code>plot = "ci"</code> and <code>line = TRUE</code> or when <code>plot = "boot"</code> and <code>line = TRUE</code> . By default, the horizontal or vertical line is drawn at 0.
<code>linetype</code>	a character string indicating the <code>linetype</code> argument in the <code>geom_hline</code> or <code>geom_vline</code> function controlling the line type of the horizontal or vertical line (default is <code>linetype = "dashed"</code>).
<code>line.col</code>	a character string indicating the <code>color</code> argument in the <code>geom_hline</code> or <code>geom_vline</code> function for controlling the color of the horizontal or vertical line.
<code>xlab</code>	a character string indicating the <code>name</code> argument in the <code>scale_x_continuous</code> function for labeling the x-axis. The default setting is <code>xlab = NULL</code> when <code>plot = "ci"</code> and <code>xlab = "Variance"</code> or <code>xlab = "Standard Deviation"</code> when <code>plot = "boot"</code> .

ylab	a character string indicating the name argument in the <code>scale_y_continuous</code> function for labeling the y-axis. The default setting is <code>ylab = "Variance"</code> or <code>ylab = "Standard Deviation"</code> when <code>plot = "ci"</code> and <code>ylab = "Probability Density, f(x)"</code> when <code>plot = "boot"</code> .
xlim	a numeric vector with two elements indicating the limits argument in the <code>scale_x_continuous</code> function for controlling the scale range of the x-axis. The default setting is <code>xlim = NULL</code> when <code>plot = "ci"</code> and <code>xlim = c(-1, 1)</code> when <code>plot = "boot"</code> .
ylim	a numeric vector with two elements indicating the limits argument in the <code>scale_y_continuous</code> function for controlling the scale range of the y-axis. The default setting is <code>ylim = c(-1, 1)</code> when <code>plot = "ci"</code> and <code>xlim = NULL</code> when <code>plot = "boot"</code> .
xbreaks	a numeric vector indicating the breaks argument in the <code>scale_x_continuous</code> function for controlling the x-axis breaks. The default setting is <code>xbreaks = NULL</code> when <code>plot = "ci"</code> and <code>xbreaks = seq(-1, 1, by = 0.25)</code> when <code>plot = "boot"</code> .
ybreaks	a numeric vector indicating the breaks argument in the <code>scale_y_continuous</code> function for controlling the y-axis breaks. The default setting is <code>ybreaks = seq(-1, 1, by = 0.25)</code> when <code>plot = "ci"</code> and <code>ybreaks = NULL</code> when <code>plot = "boot"</code> .
axis.title.size	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the axis title, i.e., <code>theme(axis.title = element_text(size = axis.text.size))</code> .
axis.text.size	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the axis text, i.e., <code>theme(axis.text = element_text(size = axis.text.size))</code> .
strip.text.size	a numeric value indicating the size argument in the <code>element_text</code> function for specifying the function controlling the font size of the strip text, i.e., <code>theme(strip.text = element_text(size = strip.text.size))</code> .
title	a character string indicating the title argument in the <code>labs</code> function for the subtitle of the plot.
subtitle	a character string indicating the subtitle argument in the <code>labs</code> function for the subtitle of the plot.
group.col	a character vector indicating the color argument in the <code>scale_color_manual</code> and <code>scale_fill_manual</code> functions when specifying a grouping variable using the argument <code>group</code> .
plot.margin	a numeric vector with four elements indicating the <code>plot.margin</code> argument in the <code>theme</code> function controlling the plot margins. The default setting is <code>c(5.5, 5.5, 5.5, 5.5)</code> , but switches to <code>c(5.5, 5.5, -2.5, 5.5)</code> when specifying a grouping variable using the argument <code>group</code> .
legend.title	a character string indicating the color argument in the <code>labs</code> function for specifying the legend title when specifying a grouping variable using the argument <code>group</code> .

<code>legend.position</code>	a character string indicating the <code>legend.position</code> in the <code>theme</code> argument for controlling the position of the legend function when specifying a grouping variable using the argument <code>group</code> . By default, the legend is placed at the bottom the plot.
<code>legend.box.margin</code>	a numeric vector with four elements indicating the <code>legend.box.margin</code> argument in the <code>theme</code> function for controlling the margins around the full legend area when specifying a grouping variable using the argument <code>group</code> .
<code>facet.ncol</code>	a numeric value indicating the <code>ncol</code> argument in the <code>facet_wrap</code> function for controlling the number of columns when specifying a split variable using the argument <code>split</code> .
<code>facet.nrow</code>	a numeric value indicating the <code>nrow</code> argument in the <code>facet_wrap</code> function for controlling the number of rows when specifying a split variable using the argument <code>split</code> .
<code>facet.scales</code>	a character string indicating the <code>scales</code> argument in the <code>facet_wrap</code> function for controlling the scales shared across facets, i.e., "fixed", "free_x", "free_y", or "free" (default) when specifying a split variable using the argument <code>split</code> .
<code>filename</code>	a character string indicating the <code>filename</code> argument including the file extension in the <code>ggsave</code> function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the <code>file</code> argument. Note that plots can only be saved when <code>plot = "ci"</code> or <code>plot = "boot"</code> .
<code>width</code>	a numeric value indicating the <code>width</code> argument (default is the size of the current graphics device) in the <code>ggsave</code> function.
<code>height</code>	a numeric value indicating the <code>height</code> argument (default is the size of the current graphics device) in the <code>ggsave</code> function.
<code>units</code>	a character string indicating the <code>units</code> argument (default is in) in the <code>ggsave</code> function.
<code>dpi</code>	a numeric value indicating the <code>dpi</code> argument (default is 600) in the <code>ggsave</code> function.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension ".txt" specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

The confidence interval based on the chi-square distribution is computed by specifying `method = "chisq"`, while the Bonett (2006) confidence interval is requested by specifying `method = "bonett"`.

By default, the Bonett confidence interval interval is computed which performs well under moderate departure from normality, while the confidence interval based on the chi-square distribution is highly sensitive to minor violations of the normality assumption and its performance does not improve with increasing sample size. Note that at least four valid observations are needed to compute the Bonett confidence interval.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	list with the input specified in <code>...</code> , <code>data</code> , <code>group</code> , and <code>split</code>
<code>args</code>	specification of function arguments
<code>boot</code>	data frame with bootstrap replicates of the variance or standard deviation when bootstrapping was requested
<code>plot</code>	ggplot2 object for plotting the results and the data frame used for plotting
<code>result</code>	result table

Note

Bootstrap confidence intervals are computed using the R package `boot` by Angelo Canty and Brian Ripley (2024).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.
- Canty, A., & Ripley, B. (2024). *boot: Bootstrap R (S-Plus) Functions*. R package version 1.3-31.
- Bonett, D. G. (2006). Approximate confidence interval for standard deviation of nonnormal distributions. *Computational Statistics and Data Analysis*, 50, 775-782. <https://doi.org/10.1016/j.csda.2004.10.003>

See Also

[ci.mean](#), [ci.mean.diff](#), [ci.median](#), [ci.prop](#), [ci.prop.diff](#), [ci.cor](#), [descript](#)

Examples

```
#-----
# Confidence Interval (CI) for the Variance

# Example 1a: Two-Sided 95% CI
ci.var(mtcars)
```

```

# Example 1b: One-Sided 99% CI based on the chi-square distributio
ci.var(mtcars, alternative = "less", method = "chisq")

#-----
# Confidence Interval (CI) for the Standard Deviation

# Example 2a: Two-Sided 95% CI
ci.sd(mtcars)

# Example 2b: One-Sided 99% CI based on the chi-square distributio
ci.sd(mtcars, alternative = "less", method = "chisq")

## Not run:
#-----
# Bootstrap Confidence Interval (CI)

# Example 3a: Bias-corrected (BC) percentile bootstrap CI
ci.var(mtcars, boot = "bc")

# Example 3b: Bias-corrected and accelerated (BCa) bootstrap CI,
# 5000 bootstrap replications, set seed of the pseudo-random number generator
ci.var(mtcars, boot = "bca", R = 5000, seed = 123)

#-----
# Grouping and Split Variable

# Example 4a: Grouping variable
ci.var(mtcars, mpg, cyl, disp, group = "vs")

# Alternative specification without using the '...' argument
ci.var(mtcars[, c("mpg", "cyl", "disp")], group = mtcars$vs)

# Example 4b: Split variable
ci.var(mtcars, mpg, cyl, disp, split = "am")

# Alternative specification without using the '...' argument
ci.var(mtcars[, c("mpg", "cyl", "disp")], split = mtcars$am)

# Example 4c: Grouping and split variable
ci.var(mtcars, mpg, cyl, disp, group = "vs", split = "am")

# Alternative specification without using the '...' argument
ci.var(mtcars[, c("mpg", "cyl", "disp")], group = mtcars$vs, split = mtcars$am)

#-----
# Write Output

# Example 5a: Text file
ci.var(mtcars, write = "CI_Var_Text.txt")

# Example 5b: Excel file
ci.var(mtcars, write = "CI_Var_Excel.xlsx")

```

```

#-----
# Plot Confidence Intervals

# Example 6a: Two-Sided 95
ci.var(mtcars, plot = "ci")

# Example 6b: Grouping variable
ci.var(mtcars, disp, hp, group = "vs", plot = "ci")

# Example 6c: Split variable
ci.var(mtcars, disp, hp, split = "am", plot = "ci")

# Example 6d: Save plot as PDF file
ci.var(mtcars, disp, hp, plot = "ci", saveplot = "CI_Var.pdf",
       width = 9, height = 6)

# Example 6e: Save plot as PNG file
ci.var(mtcars, disp, hp, plot = "ci", saveplot = "CI_Var.png",
       width = 9, height = 6)

#-----
# Plot Bootstrap Samples

# Example 7a: Two-Sided 95
ci.var(mtcars, disp, hp, boot = "bc", plot = "boot")

# Example 7b: Grouping variable
ci.var(mtcars, disp, hp, group = "vs", boot = "bc", plot = "boot")

# Example 7c: Split variable
ci.var(mtcars, disp, hp, split = "am", boot = "bc", plot = "boot")

# Example 7d: Save plot as PDF file
ci.var(mtcars, disp, hp, boot = "bc", plot = "boot",
       saveplot = "CI_Var_Boot.pdf", width = 12, height = 7)

# Example 7e: Save plot as PNG file
ci.var(mtcars, disp, hp, boot = "bc", plot = "boot",
       saveplot = "CI_Var_Boot.png", width = 12, height = 7)

## End(Not run)

```

clear

Clear Console in RStudio

Description

This function clears the console equivalent to Ctrl + L in RStudio on Windows, Mac, UNIX, or Linux operating system.

Usage

`clear()`

Author(s)

Takuya Yanagida

See Also

[restart](#), [setsource](#)

Examples

```
## Not run:

# Clear console
clear()

## End(Not run)
```

cluster.rwg	<i>Lindell, Brandt and Whitney (1999) $r^*wg(j)$ Within-Group Agreement Index for Multi-Item Scales</i>
-------------	--

Description

This function computes $r^*wg(j)$ within-group agreement index for multi-item scales as described in Lindell, Brandt and Whitney (1999).

Usage

```
cluster.rwg(data, ..., cluster, A = NULL, ranvar = NULL, z = TRUE,
            expand = TRUE, na.omit = FALSE, append = TRUE, name = "rwg",
            as.na = NULL, check = TRUE)
```

Arguments

data	a numeric vector or data frame.
...	an expression indicating the variable names in data, e.g., <code>cluster.rwg(dat, x1, x2, x3)</code> . Note that the operators <code>.</code> , <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
cluster	either a character string indicating the variable name of the cluster variable in data, or a vector representing the nested grouping structure (i.e., group or cluster variable).
A	a numeric value indicating the number of discrete response options of the items from which the random variance is computed based on $(A^2 - 1)/12$. Note that either the argument <code>j</code> or the argument <code>ranvar</code> is specified.

ranvar	a numeric value indicating the random variance to which the mean of the item variance is divided. Note that either the argument <code>j</code> or the argument <code>ranvar</code> is specified.
z	logical: if TRUE (default), Fisher z-transformation based on the formula $z = 0.5 * \log((1 + r)/(1 - r))$ is applied to the vector of <code>r*wg(j)</code> estimates.
expand	logical: if TRUE (default), vector of <code>r*wg(j)</code> estimates is expanded to match the input vector data.
na.omit	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion).
append	logical: if TRUE (default), a variable with the <code>r*wg(j)</code> within-group agreement index are appended to the data frame specified in the argument <code>data</code> .
name	a character string indicating the name of the variable appended to the data frame specified in the argument <code>data</code> when <code>append = TRUE</code> . By default, the variable is named <code>rwg</code> .
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to data, but not to <code>cluster</code> .
check	logical: if TRUE (default), argument specification is checked.

Details

The `r*wg(j)` index is calculated by dividing the mean of the item variance by the expected random variance (i.e., null distribution). The default null distribution in most research is the rectangular or uniform distribution calculated with $\sigma_e^2 u = (A^2 - 1)/12$, where A is the number of discrete response options of the items. However, what constitutes a reasonable standard for random variance is highly debated. Note that the `r*wg(j)` allows that the mean of the item variances to be larger than the expected random variances, i.e., `r*wg(j)` values can be negative.

Note that the `rwg.j.lindell()` function in the **multilevel** package uses listwise deletion by default, while the `cluster.rwg()` function uses all available information to compute the `r*wg(j)` agreement index by default. In order to obtain equivalent results in the presence of missing values, listwise deletion (`na.omit = TRUE`) needs to be applied.

Value

Returns a numeric vector containing `r*wg(j)` agreement index for multi-item scales with the same length as `group` if `expand = TRUE` or a data frame with following entries if `expand = FALSE`:

<code>cluster</code>	cluster identifier
<code>n</code>	cluster size
<code>rwg.lindell</code>	<code>r*wg(j)</code> estimate for each group
<code>z.rwg.lindell</code>	Fisher z-transformed <code>r*wg(j)</code> estimate for each cluster

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Lindell, M. K., Brandt, C. J., & Whitney, D. J. (1999). A revised index of interrater agreement for multi-item ratings of a single target. *Applied Psychological Measurement*, 23, 127-135. <https://doi.org/10.1177/01466219922031257>

O'Neill, T. A. (2017). An overview of interrater agreement on Likert scales for researchers and practitioners. *Frontiers in Psychology*, 8, Article 777. <https://doi.org/10.3389/fpsyg.2017.00777>

See Also

[cluster.scores](#)

Examples

```
dat <- data.frame(id = c(1, 2, 3, 4, 5, 6, 7, 8, 9),
                  cluster = c(1, 1, 1, 2, 2, 2, 3, 3, 3),
                  x1 = c(2, 3, 2, 1, 1, 2, 4, 3, 5),
                  x2 = c(3, 2, NA, 1, 2, 1, 3, 2, 5),
                  x3 = c(3, 1, 1, 2, 3, 3, 5, 5, 4))

# Example 1: Compute Fisher z-transformed r*wg(j) for a multi-item scale with A = 5 response options
cluster.rwg(dat, x1, x2, x3, cluster = "cluster", A = 5)

# Alternative specification without using the '...' argument
cluster.rwg(dat[, c("x1", "x2", "x3")], cluster = dat$cluster, A = 5)

# Example 2: Compute Fisher z-transformed r*wg(j) for a multi-item scale with a random variance of 2
cluster.rwg(dat, x1, x2, x3, cluster = "cluster", ranvar = 2)

# Example 3: Compute r*wg(j) for a multi-item scale with A = 5 response options
cluster.rwg(dat, x1, x2, x3, cluster = "cluster", A = 5, z = FALSE)

# Example 4: Do not expand Fisher z-transformed r*wg(j)
cluster.rwg(dat, x1, x2, x3, cluster = "cluster", A = 5, expand = FALSE)
```

cluster.scores	Cluster Scores
----------------	----------------

Description

This function computes group means by default.

Usage

```
cluster.scores(data, ..., cluster,
               fun = c("mean", "sum", "median", "var", "sd", "min", "max"),
               expand = TRUE, append = TRUE, name = ".a", as.na = NULL,
               check = TRUE)
```

Arguments

<code>data</code>	a numeric vector for centering a predictor variable, or a data frame for centering more than one predictor variable.
<code>...</code>	an expression indicating the variable names in data e.g., <code>cluster.scores(dat, x1, x2, cluster = "cluster")</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>cluster</code>	a character string indicating the variable name of the cluster variable in data, or a vector representing the nested grouping structure (i.e., group or cluster variable).
<code>fun</code>	character string indicating the function used to compute group scores, default: "mean".
<code>expand</code>	logical: if TRUE (default), vector of cluster scores is expanded to match the input vector data.
<code>append</code>	logical: if TRUE (default), cluster scores are appended to the data frame specified in the argument data.
<code>name</code>	a character string or character vector indicating the names of the computed variables. By default, variables are named with the ending ".a" resulting in e.g. "x1.a" and "x2.a". Variable names can also be specified using a character vector matching the number of variables specified in data (e.g., <code>name = c("cluster.x1", "cluster.x2")</code>).
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to the argument data, but not to <code>cluster</code> .
<code>check</code>	logical: if TRUE (default), argument specification is checked.

Value

Returns a numeric vector or data frame containing cluster scores with the same length or same number of rows as `data` if `expand = TRUE` or with the length or number of rows as `length(unique(cluster))` if `expand = FALSE`.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Hox, J., Moerbeek, M., & van de Schoot, R. (2018). *Multilevel analysis: Techniques and applications* (3rd. ed.). Routledge.
- Snijders, T. A. B., & Bosker, R. J. (2012). *Multilevel analysis: An introduction to basic and advanced multilevel modeling* (2nd ed.). Sage Publishers.

See Also

[item.scores](#), [multilevel.descript](#), [multilevel.icc](#)

Examples

```
# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

# Example 1: Compute cluster means for 'y1' and expand to match the input 'y1'
cluster.scores(Demo.twolevel, y1, cluster = "cluster", append = FALSE)

# Alternative specification without using the '...' argument
cluster.scores(Demo.twolevel$y1, cluster = Demo.twolevel$cluster)

# Example 2: Compute standard deviation for each cluster
# and expand to match the input x
cluster.scores(Demo.twolevel, cluster = "cluster", fun = "sd")

# Example 3: Compute cluster means without expanding the vector
cluster.scores(Demo.twolevel, cluster = "cluster", expand = FALSE)

# Example 4: Compute cluster means for 'y1' and 'y2' and append to 'Demo.twolevel'
cluster.scores(Demo.twolevel, y1, y2, cluster = "cluster")

# Alternative specification without using the '...' argument
cbind(Demo.twolevel,
      cluster.scores(Demo.twolevel[, c("y1", "y2")], cluster = Demo.twolevel$cluster))
```

coding

Coding Categorical Variables

Description

This function creates $k - 1$ variables for a categorical variable with k distinct levels. The coding system available in this function are dummy coding, simple coding, unweighted effect coding, weighted effect coding, repeated coding, forward Helmert coding, reverse Helmert coding, and orthogonal polynomial coding.

Usage

```
coding(data, ...,
       type = c("dummy", "simple", "effect", "weffect", "repeat",
                "fhelm", "rhelm", "poly"), base = NULL,
       name = c("dum.", "sim.", "eff.", "weff.", "rep.", "fhelm.", "rhelm.", "poly."),
       append = TRUE, as.na = NULL, check = TRUE)
```

Arguments

<code>data</code>	a numeric vector with integer values, character vector or factor.
<code>...</code>	an expression indicating the variable name in data, e.g., <code>coding(dat, x)</code> . Note that the function can only deal with one categorical variable.

type	a character string indicating the type of coding, i.e., dummy (default) for dummy coding, simple for simple coding, effect for unweighted effect coding, weffect for weighted effect coding, repeat for repeated coding, fhelm for forward Helmert coding, rhelm for reverse Helmert coding, and poly for orthogonal polynomial coding (see 'Details').
base	a numeric value or character string indicating the baseline group for dummy and simple coding and the omitted group in effect coding. By default, the first group or factor level is selected as baseline or omitted group.
name	a character string or character vector indicating the names of the coded variables. By default, variables are named "dum.", "sim.", "eff.", "weff.", "rep.", "fhelm.", "rhelm.", or "poly." depending on the type of coding with the category used in the comparison (e.g., "dum.2" and "dum.3"). Variable names can be specified using a character string (e.g., name = "dummy_" leads to dummy_2 and dummy_3) or a character vector matching the number of coded variables (e.g. name = c("x1_2", "x1_3")) which is the number of unique categories minus one.
append	logical: if TRUE (default), coded variables are appended to the data frame specified in the argument data.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
check	logical: if TRUE (default), argument specification is checked.

Details

Dummy Coding Dummy or treatment coding compares the mean of each level of the categorical variable to the mean of a baseline group. By default, the first group or factor level is selected as baseline group. The intercept in the regression model represents the mean of the baseline group. For example, dummy coding based on a categorical variable with four groups A, B, C, D makes following comparisons: B vs A, C vs A, and D vs A with A being the baseline group.

Simple Coding Simple coding compares each level of the categorical variable to the mean of a baseline level. By default, the first group or factor level is selected as baseline group. The intercept in the regression model represents the unweighted grand mean, i.e., mean of group means. For example, simple coding based on a categorical variable with four groups A, B, C, D makes following comparisons: B vs A, C vs A, and D vs A with A being the baseline group.

Unweighted Effect Coding Unweighted effect or sum coding compares the mean of a given level to the unweighted grand mean, i.e., mean of group means. By default, the first group or factor level is selected as omitted group. For example, effect coding based on a categorical variable with four groups A, B, C, D makes following comparisons: B vs (A, B, C, D), C vs (A, B, C, D), and D vs (A, B, C, D) with A being the omitted group.

Weighted Effect Coding Weighted effect or sum coding compares the mean of a given level to the weighed grand mean, i.e., sample mean. By default, the first group or factor level is selected as omitted group. For example, effect coding based on a categorical variable with four groups A, B, C, D makes following comparisons: B vs (A, B, C, D), C vs (A, B, C, D), and D vs (A, B, C, D) with A being the omitted group.

Repeated Coding Repeated or difference coding compares the mean of each level of the categorical variable to the mean of the previous adjacent level. For example, repeated coding based

on a categorical variable with four groups A, B, C, D makes following comparisons: B vs A, C vs B, and D vs C.

Forward Helmert Coding Forward Helmert coding compares the mean of each level of the categorical variable to the unweighted mean of all subsequent level(s) of the categorical variable. For example, forward Helmert coding based on a categorical variable with four groups A, B, C, D makes following comparisons: (B, C, D) vs A, (C, D) vs B, and D vs C.

Reverse Helmert Coding Reverse Helmert coding compares the mean of each level of the categorical variable to the unweighted mean of all prior level(s) of the categorical variable. For example, reverse Helmert coding based on a categorical variable with four groups A, B, C, D makes following comparisons: B vs A, C vs (A, B), and D vs (A, B, C).

Orthogonal Polynomial Coding Orthogonal polynomial coding is a form of trend analysis based on polynomials of order $k - 1$, where k is the number of levels of the categorical variable. This coding scheme assumes an ordered-categorical variable with equally spaced levels. For example, orthogonal polynomial coding based on a categorical variable with four groups A, B, C, D investigates a linear, quadratic, and cubic trends in the categorical variable.

Value

Returns a data frame with $k - 1$ coded variables or a data frame with the same length or same number of rows as ... containing the coded variables.

Note

This function uses the `contr.treatment` function from the **stats** package for dummy coding and simple coding, a modified copy of the `contr.sum` function from the **stats** package for effect coding, a modified copy of the `contr.wec` function from the **wec** package for weighted effect coding, a modified copy of the `contr.sdif` function from the **MASS** package for repeated coding, a modified copy of the `code_helmert_forward` function from the **codingMatrices** for forward Helmert coding, a modified copy of the `contr_code_helmert` function from the **faux** package for reverse Helmert coding, and the `contr.poly` function from the **stats** package for orthogonal polynomial coding.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

See Also

[rec](#), [item.reverse](#)

Examples

```
# Example 1: Dummy coding for 'gear', baseline group = 3
coding(mtcars, gear)

# Alternative specification without using the '...' argument
coding(mtcars$gear)

# Example 2: Dummy coding for 'gear', baseline group = 4
```

```

coding(mtcars, gear, base = 4)

# Example 3: Effect coding for 'gear', omitted group = 3
coding(mtcars, gear, type = "effect")

# Example 3: Effect coding for 'gear', omitted group = 4
coding(mtcars, gear, type = "effect", base = 4)

# Example 4a: Dummy-coded variable names with prefix "gear3."
coding(mtcars, gear, name = "gear3.")

# Example 4b: Dummy-coded variables named "gear_4vs3" and "gear_5vs3"
coding(mtcars, gear, name = c("gear_4vs3", "gear_5vs3"))

```

coeff.robust

Heteroscedasticity-Consistent and Cluster-Robust Standard Errors

Description

This function computes (1) heteroscedasticity-consistent or cluster-robust standard errors standard errors and significance values for (generalized) linear models estimated by using the `lm()` or the `glm()` function and (2) cluster-robust standard errors for multilevel and linear mixed-effects models estimated by using the `lmer()` function from the **lme4** package that are robust to the violation of the homoscedasticity assumption. For linear models the heteroscedasticity-robust F-test is computed as well. By default, the function uses the HC4 estimator for (generalized) linear models and the heteroscedastic-robust CR2 estimator for multilevel and linear mixed-effects models. Note that cluster-robust standard errors are available only for two-level models.

Usage

```

coeff.robust(model, cluster = NULL,
             type = c("HC0", "HC1", "HC2", "HC3", "HC4", "HC4m", "HC5",
                     "CR0", "CR1", "CR1p", "CR1S", "CR2", "CR3"),
             digits = 2, p.digits = 3, write = NULL, append = TRUE,
             check = TRUE, output = TRUE)

```

Arguments

<code>model</code>	a fitted model of class <code>lm</code> , <code>glm</code> , <code>lmerMod</code> , or <code>lmerModLmerTest</code> .
<code>cluster</code>	a vector representing the nested grouping structure (i.e., group or cluster variable). This argument is used only when requesting cluster-robust standard errors for (generalized) linear models estimated by using the <code>lm()</code> or <code>glm()</code> function. Note that the length of the vector needs to match the number of rows of the data frame used to estimate the (generalized) linear model. In the presence of missing data, the length of the vector needs to match the number of rows of the data frame after listwise deletion of missing data.

type	a character string specifying the estimation type for (generalized) linear models estimated by using the <code>lm()</code> or <code>glm()</code> function, where "H0" gives White's estimator and "H1" to "H5" are refinement of this estimator. See help page of the <code>vcovHC()</code> function in the R package <code>sandwich</code> for more details. Alternatively, a character string specifying the estimation type for multilevel and linear mixed-effects model estimated by using the <code>lmer()</code> function from the lme4 package, where "CR0" is the original form of the sandwich estimator (Liang & Zeger, 1986), "CR1" multiplies CR0 by $m/(m-1)$, where m is the number of clusters, "CR1p" multiplies CR0 by $m/(m-p)$, where p is the number of covariates, "CR1S" multiplies CR0 by $(m(N-1))/[(m-1)(N-p)]$, where N is the total number of observations, "CR2" (default) is the "bias-reduced linearization" adjustment proposed by Bell and McCaffrey (2002) and further developed in Pustejovsky and Tipton (2017), and "CR3" approximates the leave-one-cluster-out jackknife variance estimator (Bell & McCaffrey, 2002). See help page of the <code>vcovCR()</code> function in the R package <code>clubSandwich</code> for more details.
digits	an integer value indicating the number of decimal places to be used for displaying results. Note that information criteria and chi-square test statistic are printed with digits minus 1 decimal places.
p.digits	an integer value indicating the number of decimal places
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown.

Details

Heteroscedasticity-Consistent Standard Errors The family of heteroscedasticity-consistent (HC) standard errors estimator for the model parameters of a regression model is based on an HC covariance matrix of the parameter estimates and does not require the assumption of homoscedasticity. HC estimators approach the correct value with increasing sample size, even in the presence of heteroscedasticity. On the other hand, the OLS standard error estimator is biased and does not converge to the proper value when the assumption of homoscedasticity is violated (Darlington & Hayes, 2017). White (1980) introduced the idea of HC covariance matrix to econometricians and derived the asymptotically justified form of the HC covariance matrix known as HC0 (Long & Ervin, 2000). Simulation studies have shown that the HC0 estimator tends to underestimate the true variance in small to moderately large samples ($N \leq 250$) and in the presence of leverage observations, which leads to an inflated type I error risk (e.g., Cribari-Neto & Lima, 2014). The alternative estimators HC1 to HC5 are asymptotically equivalent to HC0 but include finite-sample corrections, which results in superior small sample properties compared to the HC0 estimator. Long and Ervin (2000) recommended routinely using the HC3 estimator regardless of a heteroscedasticity test. However, the HC3 estimator can be unreliable when the data contains leverage observations. The HC4 estimator,

on the other hand, performs well with small samples, in the presence of high leverage observations, and when errors are not normally distributed (Cribari-Neto, 2004). In summary, it appears that the HC4 estimator performs the best in terms of controlling the type I and type II error risk (Rosopa, 2013). As opposed to the findings of Cribari-Neto et al. (2007), the HC5 estimator did not show any substantial advantages over HC4. Both HC5 and HC4 performed similarly across all the simulation conditions considered in the study (Ng & Wilcox, 2009). Note that the F -test of significance on the multiple correlation coefficient R also assumes homoscedasticity of the errors. Violations of this assumption can result in a hypothesis test that is either liberal or conservative, depending on the form and severity of the heteroscedasticity. Hayes (2007) argued that using a HC estimator instead of assuming homoscedasticity provides researchers with more confidence in the validity and statistical power of inferential tests in regression analysis. Hence, the HC3 or HC4 estimator should be used routinely when estimating regression models. If a HC estimator is not used as the default method of standard error estimation, researchers are advised to at least double-check the results by using an HC estimator to ensure that conclusions are not compromised by heteroscedasticity. However, the presence of heteroscedasticity suggests that the data is not adequately explained by the statistical model of estimated conditional means. Unless heteroscedasticity is believed to be solely caused by measurement error associated with the predictor variable(s), it should serve as warning to the researcher regarding the adequacy of the estimated model.

Cluster-Robust Standard Errors The family of cluster-robust (CR) standard errors estimator for the model parameters of a multilevel and linear mixed-effects model are based on the heteroscedasticity-consistent (HC) standard errors estimators that have been generalized to clustered data (Zhang & Lai, 2024). The standard errors of the CR0 estimator (Liang and Zeger, 1986) rely on large samples, i.e., the CR0 estimator may result in underestimated standard errors with small number of clusters (Cameron & Miller, 2015; Imbens & Kolesar, 2016). However, there is no consensus about the minimum number of clusters, e.g., at least 100 clusters (Maas & Hox, 2004, p. 439), around 40 (Angrist & Pischke, 2008) or 30 clusters (Huang, 2016). The CR2 estimator, also referred to as Bell and McCaffrey (2002) bias-reduced linearization method, has been shown to be effective when used with a small number of clusters (Hugang & Li, 2022). For example, the CR2 estimator performed well in all conditions of a simulation study involving 20, 50, or 100 clusters regardless if homoskedasticity was violated or not. (Huang, et al, 2023). The CR3 estimator tends to over-correct the bias of the CR0 estimator, while the CR1 estimator tends to under-correct the bias (Pustejovsky & Tipton, 2018). Note that the cluster-robust SE are only robust to violation of the homoscedasticity assumption, while departure from normality or the presence of outliers can influence its performance (MacKinnon, 2012). Statistical significance testing of the regression coefficients is based on the Satterthwaite approximated degrees of freedom (Bell & McCaffrey (2002).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>model</code>	model specified in <code>model</code>
<code>args</code>	specification of function arguments
<code>result</code>	list with results, i.e., <code>coef</code> for the unstandardized regression coefficients with heteroscedasticity-consistent or cluster-robust standard errors, <code>F.test</code> for the

heteroscedasticity-robust F-Test, and sandwich for the sandwich covariance matrix

Note

The computation of heteroscedasticity-consistent standard errors is based on the `vcovHC` function from the **sandwich** package (Zeileis, Köll, & Graham, 2020) and the functions `coefTest` and `waldtest` from the `lmtest` package (Zeileis & Hothorn, 2002), while the computation of cluster-robust standard errors uses the `vcovCR` and the `coef_test` function in the **clubSandwich** package.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Angrist, J. D., & Pischke, J.-S. (2008). *Mostly harmless econometrics: An empiricist's companion*. Princeton university press.
- Bell, R. M., & McCaffrey, D. F. (2002). Bias reduction in standard errors for linear regression with multi-stage samples. *Survey Methodology*, 28(2), 169-181
- Cameron, A. C., & Miller, D. L. (2015). A practitioner's guide to cluster-robust inference. *Journal of Human Resources*, 50(2), 317-372. <https://doi.org/10.3368/jhr.50.2.317>
- Darlington, R. B., & Hayes, A. F. (2017). *Regression analysis and linear models: Concepts, applications, and implementation*. The Guilford Press.
- Cribari-Neto, F. (2004). Asymptotic inference under heteroskedasticity of unknown form. *Computational Statistics & Data Analysis*, 45, 215-233. [https://doi.org/10.1016/S0167-9473\(02\)00366-3](https://doi.org/10.1016/S0167-9473(02)00366-3)
- Cribari-Neto, F., & Lima, M. G. (2014). New heteroskedasticity-robust standard errors for the linear regression model. *Brazilian Journal of Probability and Statistics*, 28, 83-95.
- Cribari-Neto, F., Souza, T., & Vasconcellos, K. L. P. (2007). Inference under heteroskedasticity and leveraged data. *Communications in Statistics - Theory and Methods*, 36, 1877-1888. <https://doi.org/10.1080/03610920601120>
- Hayes, A.F., & Cai, L. (2007). Using heteroscedasticity-consistent standard error estimators in OLS regression: An introduction and software implementation. *Behavior Research Methods*, 39, 709-722. <https://doi.org/10.3758/BF03192961>
- Huang, F. L., & Li, X. (2022). Using cluster-robust standard errors when analyzing group-randomized trials with few clusters. *Behavior Research Methods*, 54(3), 1181-1199. <https://doi.org/10.3758/s13428-021-01627-0>
- Kuznetsova, A, Brockhoff, P. B., & Christensen, R. H. B. (2017). lmerTest Package: Tests in linear mixed effects models. *Journal of Statistical Software*, 82 13, 1-26. <https://doi.org/10.18637/jss.v082.i13>.
- Imbens, G. W., & Kolesar, M. (2016). Robust standard errors in small samples: Some practical advice. *Review of Economics and Statistics*, 98(4), 701-712. https://doi.org/10.1162/REST_a_00552
- Liang, K.-Y., & Zeger, S. L. (1986). Longitudinal data analysis using generalized linear models. *Biometrika*, 73(1), 13-22. <https://doi.org/10.1093/biomet/73.1.13>
- Long, J.S., & Ervin, L.H. (2000). Using heteroscedasticity consistent standard errors in the linear regression model. *The American Statistician*, 54, 217-224. <https://doi.org/10.1080/00031305.2000.10474549>
- Maas, C., & Hox, J. J. (2004). The influence of violations of assumptions on multilevel parameter estimates and their standard errors. *Computational Statistics & Data Analysis*, 46(3), 427-440. <https://doi.org/10.1016/j.csda.2003.08.006>

- MacKinnon, J. G. (2012). Thirty years of heteroskedasticity-robust inference. In X. Chen & N. R. Swanson (Eds.), *Recent advances and future directions in causality, prediction, and specification analysis: Essays in honor of Halbert L. White Jr* (pp. 437-461). Springer. https://doi.org/10.1007/978-1-4614-1653-1_17
- Ng, M., & Wilcoy, R. R. (2009). Level robust methods based on the least squares regression estimator. *Journal of Modern Applied Statistical Methods*, 8, 284-395. <https://doi.org/10.22237/jmasm/1257033840>
- Pustejovsky, J. E. & Tipton, E. (2018). Small sample methods for cluster-robust variance estimation and hypothesis testing in fixed effects models. *Journal of Business and Economic Statistics*, 36(4), 672-683. <https://doi.org/10.1080/07350015.2016.1247004>
- Rosopa, P. J., Schaffer, M. M., & Schroeder, A. N. (2013). Managing heteroscedasticity in general linear models. *Psychological Methods*, 18(3), 335-351. <https://doi.org/10.1037/a0032553>
- White, H. (1980). A heteroskedastic-consistent covariance matrix estimator and a direct test of heteroskedasticity. *Econometrica*, 48, 817-838. <https://doi.org/10.2307/1912934>
- Zeileis, A., & Hothorn, T. (2002). Diagnostic checking in regression relationships. *R News*, 2(3), 7-10. <http://CRAN.R-project.org/doc/Rnews/>
- Zeileis A, Köll S, & Graham N (2020). Various versatile variances: An object-oriented implementation of clustered covariances in R. *Journal of Statistical Software*, 95(1), 1-36. <https://doi.org/10.18637/jss.v095.i01>
- Zhang, Y., & Lai, M. H. C. (2024). Evaluating two small-sample corrections for fixed-effects standard errors and inferences in multilevel models with heteroscedastic, unbalanced, clustered data. *Behavior research methods*, 56(6), 5930–5946. <https://doi.org/10.3758/s13428-023-02325-9>

See Also

[coeff.std](#), [write.result](#)

Examples

```
#-----
# Example 1: Linear model

mod.lm <- lm(mpg ~ cyl + disp, data = mtcars)
coeff.robust(mod.lm)

#-----
# Example 2: Generalized linear model

mod.glm <- glm(carb ~ cyl + disp, data = mtcars, family = poisson())
coeff.robust(mod.glm)

## Not run:
#-----
# Example 3: Multilevel and Linear Mixed-Effects Model

# Load lme4 and misty package
misty::libraries(lme4, misty)

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")
```

```

# Cluster-mean centering, center() from the misty package
Demo.twolevel <- center(Demo.twolevel, x2, type = "CWC", cluster = "cluster")

# Grand-mean centering, center() from the misty package
Demo.twolevel <- center(Demo.twolevel, w1, type = "CGM", cluster = "cluster")

# Estimate two-level mixed-effects model
mod.lmer <- lmer(y1 ~ x2.c + w1.c + x2.c:w1.c + (1 + x2.c | cluster), data = Demo.twolevel)

# Statistical significance testing based on cluster-robust standard errors
coeff.robust(mod.lmer)

#-----
# Write Results

# Example 3a: Write results into a text file
coeff.robust(mod.lm, write = "Robust_Coef.txt", output = FALSE)

# Example 3b: Write results into a Excel file
coeff.robust(mod.lm, write = "Robust_Coef.xlsx", output = FALSE)

## End(Not run)

```

coeff.std	<i>Standardized Coefficients for Linear, Multilevel and Mixed-Effects Models</i>
-----------	--

Description

This function computes standardized coefficients for linear models estimated by using the `lm()` function and for multilevel and linear mixed-effects models estimated by using the `lmer()` or `lme()` function from the **lme4** or **nlme** package.

Usage

```
coeff.std(model, print = c("all", "stdx", "stdy", "stdyx"), digits = 2, p.digits = 3,
          write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>model</code>	a fitted model of class "lm", "lmerMod", "lmerModLmerTest" or "lme".
<code>print</code>	a character vector indicating which results to print, i.e. "all", for all results, "stdx" for standardizing only the predictor, "stdy" for for standardizing only the criterion, and "stdyx" for for standardizing both the predictor and the criterion. Note that the default setting is depending on the level of measurement of the predictors, i.e., if all predictors are continuous, the default setting is <code>print = "stdyx"</code> ; if all predictors are binary, the default setting is <code>print = "stdy"</code> , and if predictors are continuous and binary, the default setting is <code>print = c("stdy", "stdyx")</code> .

digits	an integer value indicating the number of decimal places to be used for displaying results.
p.digits	an integer value indicating the number of decimal places to be used for displaying the p -value.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

Linear Regression Model The linear regression model is expressed as follows:

$$y_i = \beta_0 + \beta_1 x_i + \epsilon_i$$

where y_i is the outcome variable for individual i , β_0 is the intercept, β_1 is the slope (aka regression coefficient), x_i is the predictor for individual i , and ϵ_i is the residual for individual i .

The slope β_1 estimated by using the `lm()` function can be standardized with respect to only x , only y , or both y and x :

- **StdX Standardization:** $StdX(\beta_1)$ standardizes with respect to x only and is interpreted as expected difference in y between individuals that differ one standard deviation referred to as $SD(x)$:

$$StdX(\beta_1) = \beta_1 SD(x)$$

- **StdY Standardization:** $StdY(\beta_1)$ standardizes with respect to y only and is interpreted as expected difference in y standard deviation units, referred to as $SD(y)$, between individuals that differ one unit in x :

$$StdY(\beta_1) = \frac{\beta_1}{SD(x)}$$

- **StdYX Standardization:** $StdYX(\beta_1)$ standardizes with respect to both y and x and is interpreted as expected difference in y standard deviation units between individuals that differ one standard deviation in x :

$$StdYX(\beta_1) = \beta_1 \frac{SD(x)}{SD(y)}$$

Note that the $StdYX(\beta_1)$ and the $StdY(\beta_1)$ standardizations are not suitable for the slope of a binary predictor because a one standard deviation change in a binary variable is generally not of interest (Muthen et al, 2016). Accordingly, the function does not provide the $StdYX(\beta_1)$ and the $StdY(\beta_1)$ standardizations whenever a binary vector, factor, or character vector is specified for the predictor variable.

Moderated Regression Model The moderated regression model is expressed as follows:

$$y_i = \beta_0 + \beta_1 x_{1i} + \beta_2 x_{2i} + \beta_3 x_{1i} x_{2i} + \epsilon_i$$

where β_3 is the slope for the interaction variable $x_1 x_2$.

The slope β_3 is standardized by using the product of standard deviations $SD(x_1)SD(x_2)$ rather than the standard deviation of the product $SD(x_1 x_2)$ for the interaction variable $x_1 x_2$ as discussed in Wen et al. (2010).

Note that the function does not use binary variables in the interaction term in standardizing the interaction variable. For example, when standardizing the interaction term $x_1 : x_2 : x_3$ with x_2 being binary, the product $SD(x_1)SD(x_3)$ while excluding binary predictor x_2 is used to standardize the interaction term.

Polynomial Regression Model The polynomial regression model is expressed as follows:

$$y_i = \beta_0 + \beta_1 x_i + \beta_2 x_i^2 + \epsilon_i$$

where β_2 is the slope for the quadratic term x^2 .

The slope β_2 is standardized by using the product of standard deviations $SD(x)SD(x)$ rather than the standard deviation of the product $SD(xx)$ for the quadratic term x^2 .

Multilevel and Mixed-Effects Model The random intercept and slope model in the multiple-equation notation is expressed as follows:

- Level 1:

$$y_{ij} = \beta_{0j} + \beta_{1j} x_{ij} + r_{ij}$$

- Level 2:

$$\beta_{0j} = \gamma_{00} + \gamma_{01} z_j + u_{0j}$$

$$\beta_{1j} = \gamma_{10} + u_{1j}$$

The model expressed in the single-equation notation is as follows:

$$y_{ij} = \gamma_{00} + \gamma_{10} x_{ij} + \gamma_{01} z_j + u_{0j} + u_{1j} x_{ij} + r_{ij}$$

where y_{ij} is the outcome variable for individual i in group j , γ_{00} is the fixed-effect average intercept, γ_{10} is the fixed-effect average slope for the Level-1 predictor x , and γ_{01} is the fixed-effect slope for the Level-2 predictor z .

The slopes γ_{10} and γ_{01} are standardized according to the within- and between-group or within- and between-person standard deviations, i.e., slopes are standardized with respect to the x and y standard deviation relevant for the level of the fixed effect of interest. The resulting standardized slopes are called pseudo-standardized coefficients (Hoffman 2015, p. 342). The StdYX Standardization for γ_{10} and γ_{01} is expressed as follows:

- Level-1 Predictor:

$$StdYX(\gamma_{10}) = \gamma_{10} \frac{SD(x_{ij})}{SD(y_{ij})}$$

- Level-2 Predictor:

$$StdYX(\gamma_{01}) = \gamma_{01} \frac{SD(z_j)}{SD(y_j)}$$

where $SD(x_{ij})$ and $SD(x_j)$ are the standard deviations of the predictors at each analytic level, $SD(y_{ij})$ is the square root of the Level-1 residual variance σ_r^2 and $SD(y_j)$ is square root of the Level-2 intercept variance $\sigma_{u_0}^2$ which are estimated in a null model using the `lmer` function in the **lme4** package using the restricted maximum likelihood estimation method.

The function uses the square root of the Level-1 residual variance σ_r^2 to standardize the slope of the cross-level interaction though it should be noted that it is unclear whether this is the correct approach to standardize the slope of the cross-level interaction.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame with variables used in the analysis
<code>model</code>	model specified in <code>model</code>
<code>args</code>	specification of function arguments
<code>result</code>	list with result tables, i.e., <code>coef</code> for the regression table including standardized coefficients and <code>sd</code> for the standard deviation of the outcome and predictor(s)

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Hoffman, L. (2015). *Longitudinal Analysis: Modeling Within-Person Fluctuation and Change*. Routledge.
- Muthen, B. O., Muthen, L. K., & Asparouhov, T. (2016). *Regression and mediation analysis using Mplus*. Muthen & Muthen.
- Wen, Z., Marsh, H. W., & Hau, K.-T. (2010). Structural equation models of latent interactions: An appropriate standardized solution and its scale-free properties. *Structural Equation Modeling: A Multidisciplinary Journal*, 17, 1-22. <https://doi.org/10.1080/10705510903438872>

Examples

```
#-----
# Linear Model

# Example 1a: Continuous predictors
mod.lm1 <- lm(mpg ~ cyl + disp, data = mtcars)
coeff.std(mod.lm1)

# Example 1b: Print all standardized coefficients
coeff.std(mod.lm1, print = "all")

# Example 1c: Binary predictor
mod.lm2 <- lm(mpg ~ vs, data = mtcars)
```

```

coeff.std(mod.lm2)

# Example 1d: Continuous and binary predictors
mod.lm3 <- lm(mpg ~ disp + vs, data = mtcars)
coeff.std(mod.lm3)

# Example 1e: Continuous predictors with interaction term
mod.lm4 <- lm(mpg ~ cyl*disp, data = mtcars)
coeff.std(mod.lm4)

# Example 1f: Continuous and binary predictor with interaction term
mod.lm5 <- lm(mpg ~ cyl*vs, data = mtcars)
coeff.std(mod.lm5)

# Example 1g: Continuous predictor with a quadratic term
mod.lm6 <- lm(mpg ~ cyl + I(cyl^2), data = mtcars)
coeff.std(mod.lm6)

#-----
# Multilevel and Linear Mixed-Effects Model

# Load lme4, nlme, and ggplot2 package
misty::libraries(lme4, nlme)

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

# Cluster-mean centering, center() from the misty package
Demo.twolevel <- center(Demo.twolevel, x2, type = "CWC", cluster = "cluster")

# Grand-mean centering, center() from the misty package
Demo.twolevel <- center(Demo.twolevel, w1, type = "CGM", cluster = "cluster")

# Estimate models using the lme4 package
mod1a <- lmer(y1 ~ x2.c + w1.c + (1 + x2.c | cluster), data = Demo.twolevel,
              REML = FALSE)
mod2a <- lmer(y1 ~ x2.c + w1.c + x2.c:w1.c + (1 + x2.c | cluster),
              data = Demo.twolevel, REML = FALSE)

# Estimate models using the nlme package
mod1b <- lme(y1 ~ x2.c + w1.c, random = ~ 1 + x2.c | cluster, data = Demo.twolevel,
             method = "ML")
mod2b <- lme(y1 ~ x2.c + w1.c + x2.c:w1.c, random = ~ 1 + x2.c | cluster,
             data = Demo.twolevel, method = "ML")

# Example 2: Continuous predictors
coeff.std(mod1a)
coeff.std(mod1b)

# Example 2: Continuous predictors with cross-level interaction
coeff.std(mod2a)
coeff.std(mod2b)

```

```
## Not run:
#-----
# Example 3: Write Results into a text or Excel file

# Example 3a: Text file
coeff.std(mod.lm1, write = "Std_Coef.txt", output = FALSE, check = FALSE)

# Example 3b: Excel file
coeff.std(mod.lm1, write = "Std_Coef.xlsx", output = FALSE, check = FALSE)

## End(Not run)
```

cohens.d

Cohen's d

Description

This function computes Cohen's d for one-sample, two-sample (i.e., between-subject design), and paired-sample designs (i.e., within-subject design) for one or more variables, optionally by a grouping and/or split variable. In a two-sample design, the function computes the standardized mean difference by dividing the difference between means of the two groups of observations by the weighted pooled standard deviation (i.e., Cohen's d_s according to Lakens, 2013) by default. In a paired-sample design, the function computes the standardized mean difference by dividing the mean of the difference scores by the standard deviation of the difference scores (i.e., Cohen's d_z according to Lakens, 2013) by default. Note that by default Cohen's d is computed without applying the correction factor for removing the small sample bias (i.e., Hedges' g).

Usage

```
cohens.d(x, ...)
```

Default S3 method:

```
cohens.d(x, y = NULL, mu = 0, paired = FALSE, weighted = TRUE, cor = TRUE,
  ref = NULL, correct = FALSE, alternative = c("two.sided", "less", "greater"),
  conf.level = 0.95, group = NULL, split = NULL, sort.var = FALSE,
  digits = 2, as.na = NULL, write = NULL, append = TRUE,
  check = TRUE, output = TRUE, ...)
```

S3 method for class 'formula'

```
cohens.d(formula, data, weighted = TRUE, cor = TRUE, ref = NULL,
  correct = FALSE, alternative = c("two.sided", "less", "greater"),
  conf.level = 0.95, group = NULL, split = NULL, sort.var = FALSE,
  na.omit = FALSE, digits = 2, as.na = NULL, write = NULL, append = TRUE,
  check = TRUE, output = TRUE, ...)
```

Arguments

x a numeric vector or data frame.

...	further arguments to be passed to or from methods.
y	a numeric vector.
mu	a numeric value indicating the reference mean.
paired	logical: if TRUE, Cohen's d for a paired-sample design is computed.
weighted	logical: if TRUE (default), the weighted pooled standard deviation is used to compute the standardized mean difference between two groups of a two-sample design (i.e., paired = FALSE), while standard deviation of the difference scores is used to compute the standardized mean difference in a paired-sample design (i.e., paired = TRUE).
cor	logical: if TRUE (default), paired = TRUE, and weighted = FALSE, Cohen's d for a paired-sample design while controlling for the correlation between the two sets of measurement is computed. Note that this argument is only used in a paired-sample design (i.e., paired = TRUE) when specifying weighted = FALSE.
ref	character string "x" or "y" for specifying the reference reference group when using the default cohens.d() function or a numeric value or character string indicating the reference group in a two-sample design when using the formula cohens.d() function. The standard deviation of the reference variable or reference group is used to standardized the mean difference. Note that this argument is only used in a two-sample design (i.e., paired = FALSE).
correct	logical: if TRUE, correction factor to remove positive bias in small samples is used.
alternative	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
group	a numeric vector, character vector or factor as grouping variable.
split	a numeric vector, character vector or factor as split variable.
sort.var	logical: if TRUE, output table is sorted by variables when specifying group.
digits	an integer value indicating the number of decimal places to be used for displaying results.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that as.na() function is only applied to y but not to group in a two-sample design, while as.na() function is applied to pre and post in a paired-sample design.
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.
formula	a formula of the form y ~ group for one outcome variable or cbind(y1, y2, y3) ~ group for more than one outcome variable where y is a numeric variable giving the data values and group a numeric variable, character variable or factor with two values or factor levels giving the corresponding groups.

<code>data</code>	a matrix or data frame containing the variables in the formula <code>formula</code> .
<code>na.omit</code>	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion) when specifying more than one outcome variable.

Details

Cohen (1988, p.67) proposed to compute the standardized mean difference in a two-sample design by dividing the mean difference by the unweighted pooled standard deviation (i.e., `weighted = FALSE`).

Glass et al. (1981, p. 29) suggested to use the standard deviation of the control group (e.g., `ref = 0` if the control group is coded with 0) to compute the standardized mean difference in a two-sample design (i.e., Glass's Δ) since the standard deviation of the control group is unaffected by the treatment and will therefore more closely reflect the population standard deviation.

Hedges (1981, p. 110) recommended to weight each group's standard deviation by its sample size resulting in a weighted and pooled standard deviation (i.e., `weighted = TRUE`, default). According to Hedges and Olkin (1985, p. 81), the standardized mean difference based on the weighted and pooled standard deviation has a positive small sample bias, i.e., standardized mean difference is overestimated in small samples (i.e., sample size less than 20 or less than 10 in each group). However, a correction factor can be applied to remove the small sample bias (i.e., `correct = TRUE`). Note that the function uses a gamma function for computing the correction factor, while a approximation method is used if computation based on the gamma function fails.

Note that the terminology is inconsistent because the standardized mean difference based on the weighted and pooled standard deviation is usually called Cohen's d , but sometimes called Hedges' g . Oftentimes, Cohen's d is called Hedges' d as soon as the small sample correction factor is applied. Cumming and Calin-Jageman (2017, p.171) recommended to avoid the term Hedges' g , but to report which standard deviation was used to standardized the mean difference (e.g., unweighted/weighted pooled standard deviation, or the standard deviation of the control group) and whether a small sample correction factor was applied.

As for the terminology according to Lakens (2013), in a two-sample design (i.e., `paired = FALSE`) Cohen's d_s is computed when using `weighted = TRUE` (default) and Hedges's g_s is computed when using `correct = TRUE` in addition. In a paired-sample design (i.e., `paired = TRUE`), Cohen's d_z is computed when using `weighted = TRUE`, default, while Cohen's d_{rm} is computed when using `weighted = FALSE` and `cor = TRUE`, default and Cohen's d_{av} is computed when using `weighted = FALSE` and `cor = FALSE`. Corresponding Hedges' g_z , g_{rm} , and g_{av} are computed when using `correct = TRUE` in addition.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>sample</code>	type of sample, i.e., one-, two-, or, paired-sample
<code>data</code>	matrix or data frame specified in <code>x</code>
<code>args</code>	specification of function arguments
<code>result</code>	result table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Academic Press.
- Cumming, G., & Calin-Jageman, R. (2017). *Introduction to the new statistics: Estimation, open science, & beyond*. Routledge.
- Glass, G. V., McGaw, B., & Smith, M. L. (1981). *Meta-analysis in social research*. Sage Publication.
- Goulet-Pelletier, J.-C., & Cousineau, D. (2018) A review of effect sizes and their confidence intervals, Part I: The Cohen's d family. *The Quantitative Methods for Psychology*, 14, 242-265. <https://doi.org/10.20982/tqmp.14.4.p242>
- Hedges, L. V. (1981). Distribution theory for Glass's estimator of effect size and related estimators. *Journal of Educational Statistics*, 6(3), 106-128.
- Hedges, L. V. & Olkin, I. (1985). *Statistical methods for meta-analysis*. Academic Press.
- Lakens, D. (2013). Calculating and reporting effect sizes to facilitate cumulative science: A practical primer for t-tests and ANOVAs. *Frontiers in Psychology*, 4, 1-12. <https://doi.org/10.3389/fpsyg.2013.00863>

See Also

[test.t](#), [test.z](#), [effsize](#), [cor.matrix](#), [na.auxiliary](#)

Examples

```
#-----
# One-sample design

# Example 1a: Cohen's d.z with two-sided 95% CI
# population mean = 3
cohens.d(mtcars$mpg, mu = 20)

# Example 1b: Cohen's d.z (aka Hedges' g.z) with two-sided 95% CI
# population mean = 3, with small sample correction factor
cohens.d(mtcars$mpg, mu = 20, correct = TRUE)

# Example 1c: Cohen's d.z with two-sided 95% CI
# population mean = 3, by 'vs' separately
cohens.d(mtcars$mpg, mu = 20, group = mtcars$vs)

# Example 1d: Cohen's d.z with two-sided 95% CI
# population mean = 20, split analysis by 'vs'
cohens.d(mtcars$mpg, mu = 20, split = mtcars$vs)

# Example 1e: Cohen's d.z with two-sided 95% CI
# population mean = 3, by 'vs' separately, split by 'am'
cohens.d(mtcars$mpg, mu = 20, group = mtcars$vs, split = mtcars$am)

#-----
```

```

# Two-sample design

# Example 2a: Cohen's d.s with two-sided 95% CI
# weighted pooled SD
cohens.d(mpg ~ vs, data = mtcars)

# Example 2b: Cohen's d.s with two-sided 99% CI
# weighted pooled SD
cohens.d(mpg ~ vs, data = mtcars, conf.level = 0.99)

# Example 2c: Cohen's d.s with one-sided 99% CI
# weighted pooled SD
cohens.d(mpg ~ vs, data = mtcars, alternative = "greater", conf.level = 0.99)

# Example 2d: Cohen's d.s for more than one variable with two-sided 95% CI
# weighted pooled SD
cohens.d(cbind(mpg, disp, hp) ~ vs, data = mtcars)

# Example 2e: Cohen's d with two-sided 95% CI
# unweighted SD
cohens.d(mpg ~ vs, data = mtcars, weighted = FALSE)

# Example 2f: Cohen's d.s (aka Hedges' g.s) with two-sided 95% CI
# weighted pooled SD, with small sample correction factor
cohens.d(mpg ~ vs, data = mtcars, correct = TRUE)

# Example 2g: Cohen's d (aka Hedges' g) with two-sided 95% CI
# Unweighted SD, with small sample correction factor
cohens.d(mpg ~ vs, data = mtcars, weighted = FALSE, correct = TRUE)

# Example 2h: Cohen's d (aka Glass's delta) with two-sided 95% CI
# SD of reference group 1
cohens.d(mpg ~ vs, data = mtcars, ref = 0)

# Example 2i: Cohen's d.s with two-sided 95% CI
# weighted pooled SD, by 'am' separately
cohens.d(mpg ~ vs, data = mtcars, group = mtcars$am)

# Example 2j: Cohen's d.s with two-sided 95% CI
# weighted pooled SD, split analysis by 'am'
cohens.d(mpg ~ vs, data = mtcars, split = mtcars$am)

#-----
# Paired-sample design

# Example 3a: Cohen's d.z with two-sided 95% CI
# SD of the difference scores
cohens.d(mtcars$drat, mtcars$wt, paired = TRUE)

# Example 3b: Cohen's d.z with one-sided 99% CI
# SD of the difference scores
cohens.d(mtcars$drat, mtcars$wt, paired = TRUE, alternative = "greater",
        conf.level = 0.99)

```

```

# Example 3c: Cohen's d.rm with two-sided 95% CI
# controlling for the correlation between measures
cohens.d(mtcars$drat, mtcars$wt, paired = TRUE, weighted = FALSE)

# Example 3d: Cohen's d.av with two-sided 95% CI
# without controlling for the correlation between measures
cohens.d(mtcars$drat, mtcars$wt, paired = TRUE, weighted = FALSE, cor = FALSE)

# Example 3e: Cohen's d.z (aka Hedges' g.z) with two-sided 95% CI
# SD of the difference scores
cohens.d(mtcars$drat, mtcars$wt, paired = TRUE, correct = TRUE)

# Example 3f: Cohen's d.rm (aka Hedges' g.rm) with two-sided 95% CI
# controlling for the correlation between measures
cohens.d(mtcars$drat, mtcars$wt, paired = TRUE, weighted = FALSE, correct = TRUE)

# Example 3g: Cohen's d.av (aka Hedges' g.av) with two-sided 95% CI
# without controlling for the correlation between measures
cohens.d(mtcars$drat, mtcars$wt, paired = TRUE, weighted = FALSE, cor = FALSE,
        correct = TRUE)

# Example 3h: Cohen's d.z with two-sided 95% CI
# SD of the difference scores, by 'vs' separately
cohens.d(mtcars$drat, mtcars$wt, paired = TRUE, group = mtcars$vs)

# Example 3i: Cohen's d.z with two-sided 95% CI
# SD of the difference scores, split analysis by 'vs'
cohens.d(mtcars$drat, mtcars$wt, paired = TRUE, split = mtcars$vs)

```

cor.matrix

Correlation Matrix

Description

This function computes a correlation matrix based on Pearson product-moment correlation coefficient, Spearman's rank-order correlation coefficient, Kendall's Tau-b correlation coefficient, Kendall-Stuart's Tau-c correlation coefficient, tetrachoric correlation coefficient, or polychoric correlation coefficient and computes significance values (p -values) for testing the hypothesis $H_0: \rho = 0$ for all pairs of variables.

Usage

```

cor.matrix(data, ...,
           method = c("pearson", "spearman", "kendall-b", "kendall-c", "tetra", "poly"),
           na.omit = FALSE, group = NULL, sig = FALSE, alpha = 0.05,
           print = c("all", "cor", "n", "stat", "df", "p"),
           tri = c("both", "lower", "upper"),
           p.adj = c("none", "bonferroni", "holm", "hochberg", "hommel"),

```

```

      "BH", "BY", "fdr"), continuity = TRUE,
  digits = 2, p.digits = 3, as.na = NULL,
  write = NULL, append = TRUE, check = TRUE, output = TRUE)

```

Arguments

<code>data</code>	a data frame with numeric variables, i.e., factors and character variables are excluded from data before conducting the analysis.
<code>...</code>	an expression indicating the variable names in data, e.g., <code>cor.matrix(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>method</code>	a character vector indicating which correlation coefficient is to be computed, i.e. "pearson" for Pearson product-moment correlation coefficient (default), "spearman" for Spearman's rank-order correlation coefficient, "kendall-b" for Kendall's Tau-b correlation coefficient, "kendall-c" for Kendall-Stuart's Tau-c correlation coefficient, "tetra" for tetrachoric correlation coefficient, and "poly" for polychoric correlation coefficient.
<code>na.omit</code>	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion); if FALSE (default), pairwise deletion is used.
<code>group</code>	either a character string indicating the variable name of the grouping variable in data, or a vector representing the grouping variable. Note that the grouping variable is limited to two groups.
<code>sig</code>	logical: if TRUE, statistically significant correlation coefficients are shown in boldface on the console. Note that this function does not provide statistical significance testing for tetrachoric or polychoric correlation coefficients.
<code>alpha</code>	a numeric value between 0 and 1 indicating the significance level at which correlation coefficients are printed boldface when <code>sig = TRUE</code> .
<code>print</code>	a character string or character vector indicating which results to show on the console, i.e. "all" for all results, "cor" for correlation coefficients, "n" for the sample sizes, "stat" for the test statistic, "df" for the degrees of freedom, and "p" for <i>p</i> -values. Note that the function does not provide <i>p</i> -values for tetrachoric or polychoric correlation coefficients.
<code>tri</code>	a character string indicating which triangular of the matrix to show on the console, i.e., both for upper and lower triangular, lower (default) for the lower triangular, and upper for the upper triangular.
<code>p.adj</code>	a character string indicating an adjustment method for multiple testing based on p.adjust , i.e., none, bonferroni, holm (default), hochberg, hommel, BH, BY, or fdr.
<code>continuity</code>	logical: if TRUE (default), continuity correction is used for testing Spearman's rank-order correlation coefficient and Kendall's Tau-b correlation.
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying correlation coefficients.
<code>p.digits</code>	an integer value indicating the number of decimal places to be used for displaying <i>p</i> -values.

as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

Note that unlike the `cor.test` function, this function does not compute an exact p -value for Spearman's rank-order correlation coefficient or Kendall's Tau-b correlation coefficient, but uses the asymptotic t approximation.

Statistically significant correlation coefficients can be shown in boldface on the console when specifying `sig = TRUE`. However, this option is not supported when using R Markdown, i.e., the argument `sig` will switch to FALSE.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	data frame used for the current analysis
args	specification of function arguments
result	list with result tables, i.e., <code>cor</code> for the correlation matrix, <code>n</code> for a matrix with the sample sizes, <code>stat</code> for a matrix with the test statistics, <code>df</code> for a matrix with the degrees of freedom, and <code>p-value</code> for the matrix with the significance values (p -values)

Note

This function uses the `polychoric()` function in the **psych** package by William Revelle to estimate tetrachoric and polychoric correlation coefficients.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.
- Revelle, W. (2018) *psych: Procedures for personality and psychological research*. Northwestern University, Evanston, Illinois, USA, <https://CRAN.R-project.org/package=psych> Version = 1.8.12.

See Also

[write.result](#), [cohens.d](#), [effsize](#), [multilevel.icc](#), [na.auxiliary](#), [size.cor](#).

Examples

```
# Example 1: Pearson product-moment correlation coefficient between 'Ozone' and 'Solar.R'
cor.matrix(airquality, Ozone, Solar.R)

# Alternative specification without using the '...' argument
cor.matrix(airquality[, c("Ozone", "Solar.R")])

# Example 2: Pearson product-moment correlation matrix using pairwise deletion
cor.matrix(airquality, Ozone:Wind)

# Alternative specification without using the '...' argument
cor.matrix(airquality[, c("Ozone", "Solar.R", "Wind")])

# Example 3: Spearman's rank-order correlation matrix
cor.matrix(airquality, Ozone, Solar.R, Wind, method = "spearman")

# Example 4: Pearson product-moment correlation matrix
# highlight statistically significant result at alpha = 0.05
cor.matrix(airquality, Ozone, Solar.R, Wind, sig = TRUE)

# Example 5: Pearson product-moment correlation matrix
# highlight statistically significant result at alpha = 0.05
cor.matrix(airquality, Ozone, Solar.R, Wind, sig = TRUE, alpha = 0.10)

# Example 6: Pearson product-moment correlation matrix
# print sample size and significance values
cor.matrix(airquality, Ozone, Solar.R, Wind, print = "all")

# Example 7: Pearson product-moment correlation matrix using listwise deletion,
# print sample size and significance values
cor.matrix(airquality, Ozone, Solar.R, Wind, na.omit = TRUE, print = "all")

# Example 8: Pearson product-moment correlation matrix
# print sample size and significance values with Bonferroni correction
cor.matrix(airquality, Ozone, Solar.R, Wind, na.omit = TRUE, print = "all",
           p.adj = "bonferroni")

# Example 9: Pearson product-moment correlation matrix for 'mpg', 'cyl', and 'disp'
# results for group "0" and "1" separately
cor.matrix(mtcars, mpg:disp, group = "vs")

# Alternative specification without using the '...' argument
cor.matrix(mtcars[, c("mpg", "cyl", "disp")], group = mtcars$vs)

## Not run:
# Example 10a: Write Results into a text file
cor.matrix(airquality, Ozone, Solar.R, Wind, print = "all", write = "Correlation.txt")
```

```
# Example 10b: Write Results into a Excel file
cor.matrix(airquality, Ozone, Solar.R, Wind, print = "all", write = "Correlation.xlsx")

## End(Not run)
```

crosstab

Cross Tabulation

Description

This function creates a two-way and three-way cross tabulation with absolute frequencies and row-wise, column-wise and total percentages.

Usage

```
crosstab(data, ..., print = c("no", "all", "row", "col", "total"),
         freq = TRUE, split = FALSE, na.omit = TRUE, digits = 2, as.na = NULL,
         write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a data frame with two or three columns.
<code>...</code>	an expression indicating the variable names in data, e.g., <code>crosstab(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>print</code>	a character string or character vector indicating which percentage(s) to be printed on the console, i.e., no percentages (<code>"no"</code>) (default), all percentages (<code>"all"</code>), row-wise percentages (<code>"row"</code>), column-wise percentages (<code>"col"</code>), and total percentages (<code>"total"</code>).
<code>freq</code>	logical: if TRUE (default), absolute frequencies will be included in the cross tabulation.
<code>split</code>	logical: if TRUE, output table is split in absolute frequencies and percentage(s).
<code>na.omit</code>	logical: if TRUE (default), incomplete cases are removed before conducting the analysis (i.e., listwise deletion).
<code>digits</code>	an integer indicating the number of decimal places digits to be used for displaying percentages.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension <code>".txt"</code> (e.g., <code>"Output.txt"</code>) or Excel file with file extension <code>".xlsx"</code> (e.g., <code>"Output.xlsx"</code>). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension <code>".txt"</code> specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is printed on the console.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame specified in data
<code>args</code>	specification of function arguments
<code>result</code>	list with result tables, i.e., <code>crosstab</code> for the cross tabulation, <code>freq.a</code> for the absolute frequencies, <code>perc.r</code> for the row-wise percentages, <code>perc.c</code> for the column-wise percentages, <code>perc.t</code> for the total percentages

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

[write.result](#), [freq](#), [descript](#), [multilevel.descript](#), [na.descript](#).

References

Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.

Examples

```
#-----
# Two-Dimensional Table

# Example 1: Cross Tabulation for 'vs' and 'am'
crosstab(mtcars, vs, am)

# Alternative specification without using the '...' argument
crosstab(mtcars[, c("vs", "am")])

# Example 2: Cross Tabulation, print all percentages
crosstab(mtcars, vs, am, print = "all")

# Example 3: Cross Tabulation, print row-wise percentages
crosstab(mtcars, vs, am, print = "row")

# Example 4: Cross Tabulation, print col-wise percentages
crosstab(mtcars, vs, am, print = "col")

# Example 5: Cross Tabulation, print total percentages
crosstab(mtcars, vs, am, print = "total")

# Example 6: Cross Tabulation, print all percentages, split output table
crosstab(mtcars, vs, am, print = "all", split = TRUE)

#-----
# Three-Dimensional Table
```

```
# Example 7: Cross Tabulation for 'vs', 'am', and 'gear'
crosstab(mtcars, vs:gear)

# Alternative specification without using the '...' argument
crosstab(mtcars[, c("vs", "am", "gear")])

# Example 8: Cross Tabulation, print all percentages
crosstab(mtcars, vs:gear, print = "all")

# Example 9: Cross Tabulation, print all percentages, split output table
crosstab(mtcars, vs:gear, print = "all", split = TRUE)

## Not run:
# Example 10a: Write Results into a text file
crosstab(mtcars, vs:gear, print = "all", write = "Crosstab.txt")

# Example 10b: Write Results into a Excel file
crosstab(mtcars, vs:gear, print = "all", write = "Crosstab.xlsx")

## End(Not run)
```

descript

Descriptive Statistics

Description

This function computes summary statistics for one or more than one variable, optionally by a grouping and/or split variable. By default, the function prints the number of observations (n), number of missing values (nNA), percentage of missing values (%NA), number of unique elements after omitting missing values (nUQ), arithmetic mean (M), standard deviation (SD), minimum (Min), percentage of observations at the minimum (%Min), maximum (Max), percentage of observations at the maximum (%Max), skewness (Skew), and kurtosis (Kurt).

Usage

```
descript(data, ...,
  print = c("all", "default", "n", "nNA", "pNA", "nUQ", "m", "se.m",
    "var", "sd", "min", "p.min", "p25", "med", "p75", "max", "p.max",
    "range", "iqr", "skew", "kurt"),
  group = NULL, split = NULL, sample = FALSE, sort.var = FALSE,
  na.omit = FALSE, digits = 2, as.na = NULL, write = NULL, append = TRUE,
  check = TRUE, output = TRUE)
```

Arguments

data	a numeric vector or data frame with numeric variables, i.e., factors and character variables are excluded from data before conducting the analysis.
------	---

...	an expression indicating the variable names in data, e.g., <code>descript(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
print	a character vector indicating which statistical measures to be printed on the console, i.e., <code>n</code> (number of observations), <code>nNA</code> (number of missing values), <code>pNA</code> (percentage of missing values), <code>nUQ</code> (number of unique elements after omitting missing values), <code>m</code> (arithmetic mean), <code>se.m</code> (standard error of the arithmetic mean), <code>var</code> (variance), <code>sd</code> (standard deviation), <code>med</code> (median), <code>min</code> (minimum), <code>p.min</code> (percentage of observations at the minimum), <code>p25</code> (25th percentile, first quartile), <code>p75</code> (75th percentile, third quartile), <code>max</code> (maximum), <code>p.max</code> (percentage of observations at the maximum), <code>range</code> (range), <code>iqr</code> (interquartile range), <code>skew</code> (skewness), and <code>kurt</code> (excess kurtosis). The default setting is <code>print = c("n", "nNA", "pNA", "nUQ", "m", "sd", "min", "pmin", "max", "p.max", "skew", "kurt")</code> .
group	a numeric vector, character vector or factor as grouping variable. Alternatively, a character string indicating the variable name of the grouping variable in data can be specified.
split	a numeric vector, character vector or factor as split variable. Alternatively, a character string indicating the variable name of the split variable in data can be specified.
sample	logical: if TRUE (default), the univariate sample skewness or kurtosis is computed, while the population skewness or kurtosis is computed when <code>sample = FALSE</code> .
sort.var	logical: if TRUE, output table is sorted by variables when specifying group.
na.omit	logical: if TRUE, incomplete cases are removed before conducting the analysis (i.e., listwise deletion).
digits	an integer value indicating the number of decimal places to be used.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to data, but not to group or split.
write	a character string naming a file for writing the output into either a text file with file extension <code>".txt"</code> (e.g., <code>"Output.txt"</code>) or Excel file with file extension <code>".xlsx"</code> (e.g., <code>"Output.xlsx"</code>). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension <code>.txt</code> specified in <code>write</code> , if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

Floor and Ceiling Effects This function computes the percentage of observations at both the minimum and maximum to evaluate floor and ceiling effects in continuous variables. Historically, floor or ceiling effects are considered to be present if more than 15% of observations are at the lowest or highest possible score (McHorney & Tarlov, 1995; Terwee et al., 2007). Muthen (2023, see video at 7:58) noted the rule of thumb that linear models should be avoided when the floor or ceiling effect of the outcome variable exceeds 25%.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	list with the input specified in <code>data</code> , <code>group</code> , and <code>split</code>
<code>args</code>	specification of function arguments
<code>result</code>	result table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- McHorney, C. A., & Tarlov, A. R. (1995). Individual-patient monitoring in clinical practice: are available health status surveys adequate?. *Quality of Life Research*, 4(4), 293-307. <https://doi.org/10.1007/BF01593882>
- Muthen, B. (2023, Feb. 28). *Mplus Web Talk No. 6 - Using Mplus To Do Dynamic Structural Equation Modeling: Segment 3, Descriptive Analyses* [Video]. YouTube. <https://www.statmodel.com/Webtalk6.shtml>
- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.
- Terwee, C. B., Bot, S. D., de Boer, M. R., van der Windt, D. A., Knol, D. L., Dekker, J., Bouter, L. M., & de Vet, H. C. (2007). Quality criteria were proposed for measurement properties of health status questionnaires. *Journal of Clinical Epidemiology*, 60(1), 34-42. <https://doi.org/10.1016/j.jclinepi.2006.03.012>

See Also

[ci.mean](#), [ci.mean.diff](#), [ci.median](#), [ci.prop](#), [ci.prop.diff](#), [ci.var](#), [ci.sd](#), [freq](#), [crosstab](#), [multilevel.descript](#), [na.descript](#).

Examples

```
#-----
# Descriptive statistics

# Example 1a: Descriptive statistics for 'mpg', 'cyl', and 'hp'
descript(mtcars, mpg, cyl, hp)

# Alternative specification without using the '...' argument
descript(mtcars[, c("mpg", "cyl", "hp")])

# Example 1b: Print all available statistical measures
descript(mtcars, mpg, cyl, hp, print = "all")

# Example 1c: Print default statistical measures plus median
descript(mtcars, mpg, cyl, hp, print = c("default", "med"))

#-----
```

```

# Grouping and Split Variable

# Example 2a: Grouping variable
descript(mtcars, mpg, cyl, hp, group = "vs")

# Alternative specification without using the '...' argument
descript(mtcars[, c("mpg", "cyl", "hp")], group = mtcars$vs)

# Another alternative specification without using the '...' argument
descript(mtcars[, c("mpg", "cyl", "hp", "vs")], group = "vs")

# Example 2b: Split variable
descript(mtcars, mpg, cyl, hp, split = "am")

# Alternative specification without using the '...' argument
descript(mtcars[, c("mpg", "cyl", "hp")], split = mtcars$am)

# Another alternative specification without using the '...' argument
descript(mtcars[, c("mpg", "cyl", "hp", "am")], split = "am")

# Example 2c: Grouping and split variable
descript(mtcars, mpg, cyl, hp, group = "vs", split = "am")

# Alternative specification without using the '...' argument
descript(mtcars[, c("mpg", "cyl", "hp")], group = mtcars$vs, split = mtcars$am)

# Another alternative specification without using the '...' argument
descript(mtcars[, c("mpg", "cyl", "hp", "vs", "am")], group = "vs", split = "am")

## Not run:
#-----
# Write Output

# Example 3a: Text file
descript(mtcars, write = "Descript_Text.txt")

# Example 3b: Excel file
descript(mtcars, write = "Descript_Excel.xlsx")

## End(Not run)

```

df.check

Data Check

Description

This function is a wrapper around the functions `dim` for the number of rows and columns, names for the variable names, `df.head` for the first rows, and `df.tail` for the last rows of a data frame.

Usage

```
df.check(data, print = c("dim", "names", "head", "tail"), n = 4,
         digits = 3, width = 20, row.names = TRUE, row.names.col = "gray2",
         message = TRUE, message.col = "b.blue", check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a data frame.
<code>print</code>	a character string or character vector indicating which results to show on the console, i.e., "dim", for the number of rows and number of columns, "names" for the variable names, "head" for the first rows of the data frame, and "tail" for the last rows of the data frame.
<code>n</code>	a numeric value indicating the number of rows to be printed on the console.
<code>digits</code>	a numeric value indicating the maximum number of decimal places to be used.
<code>width</code>	a numeric value indicating the maximum width of the character strings in the vector.
<code>row.names</code>	logical: if TRUE, row names of the data frame are printed on the console.
<code>row.names.col</code>	a character string indicating the text color for the row names, see color argument of the chr.color function.
<code>message</code>	logical: if TRUE, number of remaining rows and columns are printed on the console.
<code>message.col</code>	a character string indicating the text color for the number of remaining rows and columns printed on the console, see color argument of the chr.color function.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

Note that this function only provides a basic data check suitable for checking a data frame after importing data into R and is not designed to offer a thorough data check (e.g., identifying duplicate IDs or inconsistencies in the data).

Author(s)

Takuya Yanagida

See Also

[df.duplicated](#), [df.unique](#), [df.head](#), [df.tail](#), [df.long](#), [df.wide](#), [df.merge](#), [df.move](#), [df.rbind](#), [df.rename](#), [df.sort](#), [df.subset](#)

Examples

```
# Example 1: Check data frame mtcars
df.check(mtcars)
```

df.duplicated

*Extract Duplicated or Unique Rows***Description**

The function `df.duplicated` extracts duplicated rows and the function `df.unique` extracts unique rows from a matrix or data frame.

Usage

```
df.duplicated(data, ..., first = TRUE, keep.all = TRUE, from.last = FALSE,
              keep.row.names = TRUE, check = TRUE)
```

```
df.unique(data, ..., keep.all = TRUE, from.last = FALSE,
          keep.row.names = TRUE, check = TRUE)
```

Arguments

<code>data</code>	a data frame.
<code>...</code>	an expression indicating the variable names in <code>data</code> used to determine duplicated or unique rows.e.g., <code>df.duplicated(x1, x2, data = dat)</code> . Note that the operators <code>.</code> , <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see Details in the df.subset function.
<code>first</code>	logical: if TRUE (default), the <code>df.duplicated()</code> function will return duplicated rows including the first of identical rows.
<code>keep.all</code>	logical: if TRUE (default), the function will return all variables in <code>data</code> after extracting duplicated or unique rows based on the variables specified in the argument <code>...</code> .
<code>from.last</code>	logical: if TRUE, duplication will be considered from the reversed side, i.e., the last of identical rows would correspond to <code>duplicated = FALSE</code> . Note that this argument is only used when <code>first = FALSE</code> .
<code>keep.row.names</code>	logical: if TRUE (default), the row names from <code>data</code> are kept, otherwise they are set to NULL.
<code>check</code>	logical: if TRUE (default), argument specification is checked.

Details

Note that `df.unique(x)` is equivalent to `unique(x)`. That is, the main difference between the `df.unique()` and the `unique()` function is that the `df.unique()` function provides the `...` argument to specify a variable or multiple variables which are used to determine unique rows.

Value

Returns duplicated or unique rows of the data frame in `...` or `data`.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Becker, R. A., Chambers, J. M. and Wilks, A. R. (1988) *The New S Language*. Wadsworth & Brooks/Cole.

See Also

[df.check](#), [df.head](#), [df.tail](#), [df.long](#), [df.wide](#), [df.merge](#), [df.move](#), [df.rbind](#), [df.rename](#), [df.sort](#), [df.subset](#)

Examples

```
dat <- data.frame(x1 = c(1, 1, 2, 1, 4), x2 = c(1, 1, 2, 1, 6),
                 x3 = c(2, 2, 3, 2, 6), x4 = c(1, 1, 2, 2, 4),
                 x5 = c(1, 1, 4, 4, 3))

#-----
# df.duplicated() function

# Example 1: Extract duplicated rows based on all variables
df.duplicated(dat)

# Example 2: Extract duplicated rows based on 'x4'
df.duplicated(dat, x4)

# Example 3: Extract duplicated rows based on 'x2' and 'x3'
df.duplicated(dat, x2, x3)

# Example 4: Extract duplicated rows based on all variables
# exclude first of identical rows
df.duplicated(dat, first = FALSE)

# Example 5: Extract duplicated rows based on 'x2' and 'x3'
# do not return all variables
df.duplicated(dat, x2, x3, keep.all = FALSE)

# Example 6: Extract duplicated rows based on 'x4'
# consider duplication from the reversed side
df.duplicated(dat, x4, first = FALSE, from.last = TRUE)

# Example 7: Extract duplicated rows based on 'x2' and 'x3'
# set row names to NULL
df.duplicated(dat, x2, x3, keep.row.names = FALSE)

#-----
# df.unique() function

# Example 8: Extract unique rows based on all variables
df.unique(dat)
```

```
# Example 9: Extract unique rows based on 'x4'
df.unique(dat, x4)

# Example 10: Extract unique rows based on 'x1', 'x2', and 'x3'
df.unique(dat, x1, x2, x3)

# Example 11: Extract unique rows based on 'x2' and 'x3'
# do not return all variables
df.unique(dat, x2, x3, keep.all = FALSE)

# Example 12: Extract unique rows based on 'x4'
# consider duplication from the reversed side
df.unique(dat, x4, from.last = TRUE)

# Example 13: Extract unique rows based on 'x2' and 'x3'
# set row names to NULL
df.unique(dat, x2, x3, keep.row.names = FALSE)
```

df.head

*Print the First and Last Rows of a Data Frame***Description**

The function `df.head` prints the first rows of a data frame and the function `df.tail` prints the last rows of a data frame and prints as many columns as fit on the console supplemented by a summary of the remaining rows and columns.

Usage

```
df.head(data, n = 6, digits = 3, width = 20, factor.labels = TRUE,
        row.names = TRUE, row.names.col = "gray2", message = TRUE,
        message.col = "b.blue", check = TRUE, output = TRUE)

df.tail(data, n = 6, digits = 3, width = 20, factor.labels = TRUE,
        row.names = TRUE, row.names.col = "gray2", message = TRUE,
        message.col = "b.blue", check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a data frame.
<code>n</code>	a numeric value indicating the number of rows to be printed on the console.
<code>digits</code>	a numeric value indicating the maximum number of decimal places to be used.
<code>width</code>	a numeric value indicating the maximum width of the character strings in the vector.
<code>factor.labels</code>	logical: if TRUE, factor labels will be printed on the console.
<code>row.names</code>	logical: if TRUE, row names of the data frame are printed on the console.

row.names.col	a character string indicating the text color for the row names, see color argument of the chr.color function.
message	logical: if TRUE, number of remaining rows and columns are printed on the console.
message.col	a character string indicating the text color for the number of remaining rows and columns printed on the console, see color argument of the chr.color function.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Value

Returns a list with following entries:

df	data frame specified in data with the first or last n rows of the data frame with as many columns as fit on the console
row.col	character string indicating the remaining rows and columns

Author(s)

Takuya Yanagida

See Also

[df.check](#), [df.duplicated](#), [df.unique](#), [df.long](#), [df.wide](#), [df.merge](#), [df.move](#), [df.rbind](#), [df.rename](#), [df.sort](#), [df.subset](#)

Examples

```
# Example 1: Print first and last six rows
df.head(mtcars)
df.tail(mtcars)

# Example 2: Print first and last six rows without row names
df.head(mtcars, row.names = FALSE)
df.tail(mtcars, row.names = FALSE)

# Example 3: Print first and last three rows with one max. number of decimal places
df.head(mtcars, n = 3, digits = 1)
df.head(mtcars, n = 3, digits = 1)
```

df.long

*Converting Data Frames Between 'Wide' and 'Long' Format***Description**

The function `df.long` converts a data frame from the 'wide' data format (with repeated measurements in separate columns of the same row) to the 'long' data format (with repeated measurements in separate rows), while the function `df.wide` converts from the 'long' data format to the 'wide' data format.

Usage

```
df.long(data, ..., var = NULL, var.name = "value",
        time = c("num", "chr", "fac", "ord"), time.name = "time", idvar = "idvar",
        sort = TRUE, decreasing = FALSE, na.rm = FALSE, check = TRUE)

df.wide(data, ..., var, var.name = var, time = "time", idvar = "idvar",
        sep = "", check = TRUE)
```

Arguments

<code>data</code>	a data frame in 'wide' or 'long' format.
<code>...</code>	an expression indicating the time-invariant variable names in data that should be kept after converting data to the 'long' or 'wide' format. Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function. Note that the <code>...</code> is not specified when all variables should be kept in the converted data frame.
<code>var</code>	a character vector (one set of variable names) or a list of character vectors (multiple sets of variables names) in the wide data format indicating the sets of time-varying variables in the wide format that correspond to single variables in the long format when using the <code>df.long</code> function. Note that all variables excluded those specified in the argument <code>...</code> are used when <code>var = NULL</code> (default), see Example 7. A character vector indicating the variable name(s) in the long format that are being split into separate variables when using the <code>df.wide</code> function.
<code>var.name</code>	a character vector specifying the variable names in the long format that correspond to the sets of time-varying variables in the wide data format when using the <code>df.long</code> function or a character vector specifying the prefix of the variable names in the wide format that correspond to the time-varying variables in the long format.
<code>time</code>	a character string indicating the data type of the newly created variable in the long format when using the <code>df.long</code> function, i.e., "num" for numeric consecutive integers starting from 0 (e.g., 0, 1, 2, 3 for a set of four variables in the wide data format), "chr" for a character vector, "fac" for a factor, and "ord" for an ordered factor. Note that the variable names of the set of variables in the wide data format is used when specifying "chr", "fac", or "ord" if only one

set of variables is specified in the "var" argument. Otherwise numeric consecutive integers starting from 1 as character, factor or ordered factor are used. Or a character string indicating the variable name in the long data format that differentiates multiple records from the same group or individual when using the `df.wide` function.

<code>time.name</code>	a character string indicating the name of the newly created variable in the long format when using the <code>df.long</code> function. By default, the variable is named "time". Note that variable names can also be specified using the <code>var</code> when multiple sets of time-varying variables are specified in a list, e.g., <code>var = list(dep = c("ad", "bd"), anx = c("aa", "ba"))</code> (see alternative specification in Example 5).
<code>idvar</code>	a character string indicating the name of the identification variable in the wide data format that is used to sort the data after converting a data frame from wide to long format when using the <code>df.long</code> function and specifying <code>sort = TRUE</code> . Note that the function will create an identification variable with consecutive integer starting from 1 if the variable specified in <code>idvar</code> is not found in data. Or a character string indicating the name of the identification variable in the long data format when using the <code>df.wide</code> function.
<code>sort</code>	logical: if TRUE (default), data frame in the long format is sorted according to the identification variable specified in <code>idvar</code> when using the <code>df.long</code> function.
<code>decreasing</code>	logical: if TRUE, the sort is decreasing when specifying <code>sort = TRUE</code> .
<code>na.rm</code>	logical: if TRUE, rows with NA values for all variables in the long format that correspond to the sets of time-varying variables in the wide data format will be removed from the data when using the <code>df.long</code> function.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>sep</code>	a character string indicating a separating character in the variable names after converting data from the long format to the wide format when using the <code>df.wide</code> function. For example, the variable value in the long format will be split into the variables <code>value0</code> , <code>value1</code> , and <code>value2</code> when specifying <code>sep = ""</code> (default), but will be split into the variables <code>value_0</code> , <code>value_1</code> , and <code>value_2</code> when specifying <code>sep = "_"</code> .

Value

Data frame that is converted to the 'long' or 'wide' format.

Note

The function `df.long` uses the function `melt` and the function `df.long` uses the function `dcast` provided in the R package **data.table** by Tyson Barrett et al., (2025).

Author(s)

Takuya Yanagida

References

Barrett, T., Dowle, M., Srinivasan, A., Gorecki, J., Chirico, M., Hocking, T., & Schwendinger, B. (2025). data.table: Extension of 'data.frame'. R package version 1.17.8. <https://CRAN.R-project.org/package=data.table>

See Also

[df.check](#), [df.duplicated](#), [df.unique](#), [df.head](#), [df.tail](#), [df.merge](#), [df.move](#), [df.rbind](#), [df.rename](#), [df.sort](#), [df.subset](#)

Examples

```
dat.w <- data.frame(id = c(23, 55, 71),
                    gend = c("male", "female", "male"), age = c(22, 19, 26),
                    adep = c(3, 6, NA), bdep = c(5, 5, 6), cdep = c(4, NA, 5),
                    aanx = c(5, 3, 6), banx = c(NA, 7, 2), canx = c(6, NA, 8))

#-----
# Convert from 'wide' data format to the 'long' data format

# Example 1: One set of time-varying variables combined into "dep"
df.long(dat.w, var = c("adep", "bdep", "cdep"), var.name = "dep", idvar = "id")

# Example 2: Select time-invariant variables 'gend' and 'age'
df.long(dat.w, gend, age, var = c("adep", "bdep", "cdep"), var.name = "dep",
        idvar = "id")

# Example 3: Newly created variable "type" as character vector
df.long(dat.w, age, var = c("adep", "bdep", "cdep"), var.name = "dep",
        idvar = "id", time = "chr", time.name = "type")

# Example 4: User-defined variable "type"
df.long(dat.w, age, var = c("adep", "bdep", "cdep"), var.name = "dep",
        idvar = "id", time = c("pre", "post", "follow-up"), time.name = "type")

# Example 5: Two sets of time-varying variables combined into "dep" and "anx"
df.long(dat.w, age,
        var = list(c("adep", "bdep", "cdep"), c("aanx", "banx", "canx")),
        var.name = c("dep", "anx"), idvar = "id")

# Alternative specification using named lists for the argument 'var'
df.long(dat.w, age,
        var = list(dep = c("adep", "bdep", "cdep"), anx = c("aanx", "banx", "canx")),
        idvar = "id")

# Example 6: Remove rows with only NA values
df.long(dat.w, age, var = list(c("adep", "bdep", "cdep"), c("aanx", "banx", "canx")),
        idvar = "id", sort = FALSE, na.rm = TRUE)

# Example 7: Convert all variables except "age" and "gend"
df.long(dat.w, age, gend, idvar = "id")
```

```
#-----
# Convert from 'long' data format to the 'wide' data format

dat.l <- df.long(dat.w,
  var = list(c("adep", "bdep", "cdep"), c("aanx", "banx", "canx")),
  var.name = c("dep", "anx"), idvar = "id")

# Example 8: Time-varying variables "dep" and "anx" expanded into multiple variables
df.wide(dat.l, var = c("dep", "anx"), idvar = "id", time = "time")

# Example 9: Select time-invariant variables 'age'
df.wide(dat.l, age, var = c("dep", "anx"), idvar = "id", time = "time")

# Example 10: Variable name prefix of the expanded variables "depre" and "anxie"
#           with separating character "."
df.wide(dat.l, var = c("dep", "anx"), var.name = c("depre", "anxie"),
  idvar = "id", time = "time", sep = ".")
```

df.merge

*Merge Multiple Data Frames***Description**

This function merges data frames by a common column (i.e., matching variable).

Usage

```
df.merge(..., by, all = TRUE, check = TRUE, output = TRUE)
```

Arguments

...	a sequence of matrices or data frames and/or matrices to be merged to one.
by	a character string indicating the column used for merging (i.e., matching variable), see 'Details'.
all	logical: if TRUE (default), then extra rows with NAs will be added to the output for each row in a data frame that has no matching row in another data frame.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

There are following requirements for merging multiple data frames: First, each data frame has the same matching variable specified in the `by` argument. Second, matching variable in the data frames have all the same class. Third, there are no duplicated values in the matching variable in each data frame. Fourth, there are no missing values in the matching variables. Last, there are no duplicated variable names across the data frames except for the matching variable.

Note that it is possible to specify data frames matrices and/or in the argument ... However, the function always returns a data frame.

Value

Returns a merged data frame.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

See Also

[df.check](#), [df.duplicated](#), [df.unique](#), [df.head](#), [df.tail](#), [df.long](#), [df.wide](#), [df.move](#), [df.rbind](#), [df.rename](#), [df.sort](#), [df.subset](#)

Examples

```
adat <- data.frame(id = c(1, 2, 3),
                  x1 = c(7, 3, 8))

bdat <- data.frame(id = c(1, 2),
                  x2 = c(5, 1))

cdat <- data.frame(id = c(2, 3),
                  y3 = c(7, 9))

ddat <- data.frame(id = 4,
                  y4 = 6)

# Example 1: Merge 'adat', 'bdat', 'cdat', and 'ddat' by the variable 'id'
df.merge(adat, bdat, cdat, ddat, by = "id")

# Example 2: Do not show output on the console
df.merge(adat, bdat, cdat, ddat, by = "id", output = FALSE)

## Not run:
#-----
# Error messages

adat <- data.frame(id = c(1, 2, 3),
                  x1 = c(7, 3, 8))

bdat <- data.frame(code = c(1, 2, 3),
                  x2 = c(5, 1, 3))

cdat <- data.frame(id = factor(c(1, 2, 3)),
                  x3 = c(5, 1, 3))

ddat <- data.frame(id = c(1, 2, 2),
                  x2 = c(5, 1, 3))

edat <- data.frame(id = c(1, NA, 3),
                  x2 = c(5, 1, 3))

fdat <- data.frame(id = c(1, 2, 3),
```

```

x1 = c(5, 1, 3))

# Error 1: Data frames do not have the same matching variable specified in 'by'.
df.merge(adat, bdat, by = "id")

# Error 2: Matching variable in the data frames do not all have the same class.
df.merge(adat, cdat, by = "id")

# Error 3: There are duplicated values in the matching variable specified in 'by'.
df.merge(adat, ddat, by = "id")

# Error 4: There are missing values in the matching variable specified in 'by'.
df.merge(adat, edat, by = "id")

# Error 5: There are duplicated variable names across data frames.
df.merge(adat, fdat, by = "id")

## End(Not run)

```

df.move

*Move Variable(s) in a Data Frame***Description**

This function moves variables to a different position in the data frame, i.e., changes the column positions in the data frame. By default, variables specified in the first argument ... are moved to the first position in the data frame specified in the argument data.

Usage

```
df.move(data, ..., before = NULL, after = NULL, first = TRUE, check = TRUE)
```

Arguments

data	a data frame.
...	an expression indicating the variable names in data to move. Note that the operators +, -, ~, :, :~, and ! can also be used to select variables, see Details in the df.subset function.
before	a character string indicating a variable in data. Variable(s) specified in ... are moved to the left-hand side of this variable.
after	a character string indicating a variable in data. Variable(s) specified in ... are moved to the right-hand side of this variable.
first	logical: if TRUE (default), variable(s) specified in ... will be moved to the first position in 'data', if FALSE, variable(s) specified in ... will be moved to the last position in 'data'.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns the data frame in data with columns in a different place.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Becker, R. A., Chambers, J. M. and Wilks, A. R. (1988) *The New S Language*. Wadsworth & Brooks/Cole.

See Also

[df.check](#), [df.duplicated](#), [df.unique](#), [df.head](#), [df.tail](#), [df.long](#), [df.wide](#), [df.merge](#), [df.rbind](#), [df.rename](#), [df.sort](#), [df.subset](#)

Examples

```
# Example 1: Move variables 'hp' and 'am' to the first position
df.move(mtcars, hp, am)

# Example 2: Move variables 'hp' and 'am' to the last position
df.move(mtcars, hp, am, first = FALSE)

# Example 3: Move variables 'hp' and 'am' to the left-hand side of 'disp'
df.move(mtcars, hp, am, before = "disp")

# Example 4: Move variables 'hp' and 'am' to the right-hand side of 'disp'
df.move(mtcars, hp, am, after = "disp")
```

df.rbind

Combine Data Frames by Rows, Filling in Missing Columns

Description

This function takes a sequence of data frames and combines them by rows, while filling in missing columns with NAs.

Usage

```
df.rbind(...)
```

Arguments

... a sequence of data frame to be row bind together. This argument can be a list of data frames, in which case all other arguments are ignored. Any NULL inputs are silently dropped. If all inputs are NULL, the output is also NULL.

Details

This is an enhancement to `rbind` that adds in columns that are not present in all inputs, accepts a sequence of data frames, and operates substantially faster.

Column names and types in the output will appear in the order in which they were encountered.

Unordered factor columns will have their levels unified and character data bound with factors will be converted to character. POSIXct data will be converted to be in the same time zone. Array and matrix columns must have identical dimensions after the row count. Aside from these there are no general checks that each column is of consistent data type.

Value

Returns a single data frame

Note

This function is a copy of the `rbind.fill()` function in the **plyr** package by Hadley Wickham.

Author(s)

Hadley Wickham

References

Wickham, H. (2011). The split-apply-combine strategy for data analysis. *Journal of Statistical Software*, 40, 1-29. <https://doi.org/10.18637/jss.v040.i01>

Wickham, H. (2019). *plyr: Tools for Splitting, Applying and Combining Data*. R package version 1.8.5.

See Also

`df.check`, `df.duplicated`, `df.unique`, `df.head`, `df.tail`, `df.long`, `df.wide`, `df.merge`, `df.move`, `df.rename`, `df.sort`, `df.subset`

Examples

```
adat <- data.frame(id = c(1, 2, 3), a = c(7, 3, 8), b = c(4, 2, 7))
bdat <- data.frame(id = c(4, 5, 6), a = c(2, 4, 6), c = c(4, 2, 7))
cdat <- data.frame(id = c(7, 8, 9), a = c(1, 4, 6), d = c(9, 5, 4))

# Example 1
df.rbind(adat, bdat, cdat)
```

df.rename

*Rename Columns in a Matrix or Variables in a Data Frame***Description**

This function renames columns in a matrix or variables in a data frame by (1) using `old_name = new_name`, by using the functions `toupper`, `tolower`, `sub`, and `gsub`, or (3) by specifying a character vector indicating the column(s) or variable(s) to be renamed (argument `from`) and a character vector indicating the corresponding replacement values (argument `to`).

Usage

```
df.rename(data, ..., from, to, check = TRUE)
```

Arguments

<code>data</code>	a matrix or data frame.
<code>...</code>	<code>old_name = new_name</code> when <code>from = NULL</code> and <code>to = NULL</code> , or one of the functions <code>toupper</code> , <code>tolower</code> , <code>sub</code> , and <code>gsub</code> . Note that a tilde (<code>~</code>) needs to be specified before when using a function, e.g., <code>~toupper</code> or <code>~gsub("_", ".")</code> .
<code>from</code>	a character string or character vector indicating the column(s) or variable(s) to be renamed.
<code>to</code>	a character string or character vector indicating the corresponding replacement values for the column(s) or variable(s) specified in the argument <code>name</code> .
<code>check</code>	logical: if <code>TRUE</code> (default), argument specification is checked.

Value

Returns the matrix or data frame `data` with renamed columns or variables.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

See Also

[df.check](#), [df.duplicated](#), [df.unique](#), [df.head](#), [df.tail](#), [df.long](#), [df.wide](#), [df.merge](#), [df.move](#), [df.rbind](#), [df.sort](#), [df.subset](#)

Examples

```
#-----
# Rename using variable names

# Example 1a: Rename 'cyl' in 'mtcars' to 'cylinder' using 'old_name = new_name'
df.rename(mtcars, cyl = cylinder)
```

```
# Example 1b: Rename 'cyl' in 'mtcars' to 'cylinder' using 'from' and 'to'
df.rename(mtcars, from = "cyl", to = "cylinder")

# Example 2a: Rename 'cyl' and 'wt' in 'mtcars' to 'cylinder' and 'weight'
# using 'old_name = new_name'
df.rename(mtcars, cyl = cylinder, wt = weight)

# Example 2b: Rename 'cyl' and 'wt' in 'mtcars' to 'cylinder' and 'weight'
# using using 'from' and 'to'
df.rename(mtcars, from = c("cyl", "wt"), to = c("cylinder", "weight"))

#-----
# Rename using functions

# Example 3: Convert all variable names to lowercase
df.rename(iris, ~tolower)

# Example 4: Replace all '.' with '_'
# Note, the argument fixed is set to TRUE by default.
df.rename(iris, ~gsub(".", "_"))

# Example 5: Replace all 'S' with 'P'
df.rename(iris, ~gsub("S", "P"))

# Example 6: Replace all 'S' with 'P', ignore case during matching
df.rename(iris, ~gsub("S", "P", ignore.case = TRUE))
```

df.sort	<i>Data Frame Sorting</i>
---------	---------------------------

Description

This function arranges a data frame in increasing or decreasing order according to one or more variables.

Usage

```
df.sort(data, ..., decreasing = FALSE, check = TRUE)
```

Arguments

- data a data frame.
- ... a sorting variable or a sequence of sorting variables which are specified without quotes ' ' or double quotes "".
- decreasing logical: if TRUE, the sort is decreasing.
- check logical: if TRUE (default), argument specification is checked.

Value

Returns data frame data sorted according to the variables specified in . . . , a matrix will be coerced to a data frame.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Becker, R. A., Chambers, J. M. and Wilks, A. R. (1988) *The New S Language*. Wadsworth & Brooks/Cole.

Knuth, D. E. (1998) *The Art of Computer Programming, Volume 3: Sorting and Searching* (2nd ed.). Addison-Wesley.

See Also

[df.check](#), [df.duplicated](#), [df.unique](#), [df.head](#), [df.tail](#), [df.long](#), [df.wide](#), [df.merge](#), [df.move](#), [df.rbind](#), [df.rename](#), [df.subset](#)

Examples

```
# Example 1: Sort data frame 'mtcars' by 'mpg' in increasing order
df.sort(mtcars, mpg)

# Example 2: Sort data frame 'mtcars' by 'mpg' in decreasing order
df.sort(mtcars, mpg, decreasing = TRUE)

# Example 3: Sort data frame 'mtcars' by 'mpg' and 'cyl' in increasing order
df.sort(mtcars, mpg, cyl)

# Example 4: Sort data frame 'mtcars' by 'mpg' and 'cyl' in decreasing order
df.sort(mtcars, mpg, cyl, decreasing = TRUE)
```

`df.subset`*Subsetting Data Frames*

Description

This function returns subsets of data frames which meet conditions.

Usage

```
df.subset(data, ..., subset = NULL, drop = TRUE, check = TRUE)
```

Arguments

data	a data frame.
...	an expression indicating variables to select from the data frame specified in data. See Details for the list of operators used in this function, i.e., +, -, ~, :, ::, and !. Note that all variables are selected if the argument ... is not specified.
subset	a logical expression indicating rows to keep, e.g., var == 1, var1 == 1 & var2 == 3, or gender == "female". By default, all rows of the data frame specified in data are kept. Note that logical queries for rows resulting in missing values are not select.
drop	logical: if TRUE (default), data frame with a single column is converted into a vector.
check	logical: if TRUE (default), argument specification is checked.

Details

The argument ... is used to specify an expression indicating the variables to select and/or remove from the data frame specified in data. There are six operators which can be used in the expression ...:

Plus (+) Operator The plus operator is used to select variables matching a prefix from the data frame specified in data. For example, `df.subset(dat, +x)` selects all variables with the prefix x. Note that this operator is equivalent to the function `starts_with()` from the **tidyselect** package.

Minus (-) Operator The minus operator is used to select variables matching a suffix from the data frame specified in data. For example, `df.subset(dat, -y)` selects all variables with the suffix y. Note that this operator is equivalent to the function `ends_with()` from the **tidyselect** package.

Tilde (~) Operator The tilde operator is used to select variables containing a word from the data frame specified in data. For example, `df.subset(dat, ~a1)` selects all variables with the word a1. Note that this operator is equivalent to the function `contains()` from the **tidyselect** package.

Colon (:) operator The colon operator is used to select a range of consecutive variables from the data frame specified in data. For example, `df.subset(dat, x:z)` selects all variables from x to z. Note that this operator is equivalent to the `:` operator from the `select` function in the **dplyr** package.

Double Colon (: :) Operator The double colon operator is used to select numbered variables from the data frame specified in data. For example, `df.subset(dat, x1:x3)` selects the variables x1, x2, and x3. Note that this operator is similar to the function `num_range()` from the **tidyselect** package.

Exclamation Point (!) Operator The exclamation point operator is used to drop variables from the data frame specified in the argument data or for taking the complement of a set of variables. For example, `df.subset(dat, !x)` selects all variables except the variable x, `df.subset(dat, !~x)` selects all variables except variables with the prefix x, or `df.subset(dat, x1:x10, !x3:x5)` selects all variables from x1 to x10 but excludes all variables from x3 to x5. Note that this operator is equivalent to the `!` operator from the `select` function in the **dplyr** package.

Operators can be combined within the same function call. For example, `df.subset(dat, +x, -y, !x2:x4, z)` selects all variables with the prefix `x` and with the suffix `y` but excludes variables from `x2` to `x4` and select variable `z`.

Value

Returns a data frame containing the variables and rows selected in the argument `...` and rows selected in the argument `subset`.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Becker, R. A., Chambers, J. M. and Wilks, A. R. (1988) *The New S Language*. Wadsworth & Brooks/Cole.

See Also

[df.check](#), [df.duplicated](#), [df.unique](#), [df.head](#), [df.tail](#), [df.long](#), [df.wide](#), [df.merge](#), [df.move](#), [df.rbind](#), [df.rename](#), [df.sort](#),

Examples

```
## Not run:
#-----
# Select single variables

# Example 1: Select 'Sepal.Length' and 'Petal.Width'
df.subset(iris, Sepal.Length, Petal.Width)

#-----
# Select rows

# Example 2a: Select all variables, select rows with 'Species' equal 'setosa'
df.subset(iris, subset = Species == "setosa")

# Example 2b: Select all variables, select rows with 'Petal.Length' smaller 1.2
df.subset(iris, subset = Petal.Length < 1.2)

#-----
# Select variables matching a prefix using the + operator

# Example 3: Select variables with prefix 'Petal'
df.subset(iris, +Petal)

#-----
# Select variables matching a suffix using the - operator

# Example 4: Select variables with suffix 'Width'
df.subset(iris, -Width)
```

```

#-----
# Select variables containing a word using the ~ operator
#
# Example 5: Select variables containing 'al'
df.subset(iris, ~al)

#-----
# Select consecutive variables using the : operator

# Example 6: Select all variables from 'Sepal.Width' to 'Petal.Width'
df.subset(iris, Sepal.Width:Petal.Width)

#-----
# Select numbered variables using the :: operator

# Example 7: Select all variables from 'x1' to 'x3' and 'y1' to 'y3'
df.subset(anscombe, x1::x3, y1::y3)
#
#-----
# Drop variables using the ! operator

# Example 8a: Select all variables except 'Sepal.Width'
df.subset(iris, !Sepal.Width)

# Example 8b: Select all variables except variables with prefix 'Petal'
df.subset(iris, !+Petal)

# Example 8c: Select all variables except variables with suffix 'Width'
df.subset(iris, !-Width)

# Example 8d: Select all variables except 'Sepal.Width' to 'Petal.Width'
df.subset(iris, !Sepal.Width:Petal.Width)

#-----
# Combine +, -, !, and : operators

# Example 9: Select variables with prefix 'x' and suffix '3', but exclude
# variables from 'x2' to 'x3'
df.subset(anscombe, +x, -3, !x2:x3)

## End(Not run)

```

dominance

Dominance Analysis

Description

This function conducts dominance analysis (Budescu, 1993; Azen & Budescu, 2003) for linear models estimated by using the `lm()` function to determine the relative importance of predictor vari-

ables. By default, the function reports general dominance, but conditional and complete dominance can be requested by specifying the argument `print`.

Usage

```
dominance(model, print = c("all", "gen", "cond", "comp"), digits = 3,
          write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>model</code>	a fitted model of class <code>lm</code> .
<code>print</code>	a character string or character vector indicating which results to show on the console, i.e. "all" for all results, "gen" for general dominance, "cond" for conditional dominance, and "comp" for complete dominance.
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying results. Note that the percentage relative importance of predictors are printed with <code>digits</code> minus 1 decimal places.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension ".txt" specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown.

Details

Dominance analysis (Budescu, 1993; Azen & Budescu, 2003) is used to determine the relative importance of predictor variables in a statistical model by examining the additional contribution of predictors in *R*-squared relative to each other in all of the possible $2^{(p-2)}$ subset models with p being the number of predictors. Three levels of dominance can be established through pairwise comparison of all predictors in a regression model:

Complete Dominance A predictor completely dominates another predictor if its additional contribution in *R*-Squared is higher than that of the other predictor across all possible subset models that do not include both predictors. For example, in a regression model with four predictors, X_1 completely dominates X_2 if the additional contribution in *R*-squared for X_1 is higher compared to X_2 in (1) the null model without any predictors, (2) the model including X_3 , (3) the model including X_4 , and (4) the model including both X_3 and X_4 . Note that complete dominance cannot be established if one predictor's additional contribution is greater than the other's for some, but not all of the subset models. In this case, dominance is undetermined and the result will be NA

Conditional Dominance A predictor conditionally dominates another predictor if its average additional contribution in *R*-squared is higher within each model size than that of the other predictor. For example, in a regression model with four predictors, X_1 conditionally dominates X_2 if the average additional contribution in *R*-squared is higher compared to X_2 in (1) the null

model without any predictors, (2) the four models including one predictor, (3) the six models including two predictors, and (4) the four models including three predictors.

General Dominance A predictor generally dominates another predictor if its overall averaged additional contribution in R -squared is higher than that of the other predictor. For example, in a regression model with four predictors, X_1 generally dominates X_2 if the average across the four conditional values (i.e., null model, model with one predictor, model with two predictors, and model with three predictors) is higher than that of X_2 . Note that the general dominance measures represent the proportional contribution that each predictor makes to the R -squared since their sum across all predictors equals the R -squared of the full model.

The three levels of dominance are related to each other in a hierarchical fashion: Complete dominance implies conditional dominance, which in turn implies general dominance. However, the converse may not hold for more than three predictors. That is, general dominance does not imply conditional dominance, and conditional dominance does not necessarily imply complete dominance.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>model</code>	model specified in <code>model</code>
<code>args</code>	specification of function arguments
<code>result</code>	list with results, i.e., <code>gen</code> for general dominance, <code>cond</code> for conditional dominance, <code>comp</code> for complete dominance, and <code>condtsat</code> for the statistics of the conditional dominance

Note

This function is based on the `domir` function from the `domir` package (Luchman, 2023).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Azen, R., & Budescu, D. V. (2003). The dominance analysis approach for comparing predictors in multiple regression. *Psychological Methods*, 8(2), 129–148. <https://doi.org/10.1037/1082-989X.8.2.129>
- Budescu, D. V. (1993). Dominance analysis: A new approach to the problem of relative importance of predictors in multiple regression. *Psychological Bulletin*, 114(3), 542–551. <https://doi.org/10.1037/0033-2909.114.3.542>
- Luchman J (2023). *domir: Tools to support relative importance analysis*. R package version 1.0.1, <https://CRAN.R-project.org/package=domir>.

See Also

[dominance.manual](#), [coeff.std](#), [write.result](#)

Examples

```
#-----
# Example 1: Dominance analysis for a linear model

# Example 1
mod <- lm(mpg ~ cyl + disp + hp, data = mtcars)
dominance(mod)

# Print all results
dominance(mod, print = "all")

## Not run:
#-----
# Write results into a Text or Excel file

# Example 2a: Text file
dominance(mod, write = "Dominance.txt", output = FALSE)

# Example 2b: Excel file
dominance(mod, write = "Dominance.xlsx", output = FALSE)
## End(Not run)
```

dominance.manual

Dominance Analysis, Manually Inputting a Correlation Matrix

Description

This function conducts dominance analysis (Budescu, 1993; Azen & Budescu, 2003) based on a (model-implied) correlation matrix of the manifest or latent variables. Note that the function only provides general dominance.

Usage

```
dominance.manual(x, out = NULL, digits = 3, write = NULL, append = TRUE,
                 check = TRUE, output = TRUE)
```

Arguments

x	a matrix or data frame with the (model-implied) correlation matrix of the manifest or latent variables. Note that column names need to represent the variables names in x.
out	a character string representing the outcome variable. By default, the first row and column represents the outcome variable.

<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying results. Note that the percentage relative importance of predictors are printed with <code>digits</code> minus 1 decimal places.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>x</code>	correlation matrix specified in <code>x</code>
<code>args</code>	specification of function arguments
<code>result</code>	results table for the general dominance

Note

This function implements the function provided in Appendix 1 of Gu (2022) and copied the function `combinations()` from the `gtools` package (Bolker, Warnes, & Lumley, 2022).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Azen, R., & Budescu, D. V. (2003). The dominance analysis approach for comparing predictors in multiple regression. *Psychological Methods*, 8(2), 129–148. <https://doi.org/10.1037/1082-989X.8.2.129>
- Bolker, B., Warnes, G., & Lumley, T. (2022). *gtools: Various R Programming Tools*. R package version 3.9.4, <https://CRAN.R-project.org/package=gtools>
- Budescu, D. V. (1993). Dominance analysis: A new approach to the problem of relative importance of predictors in multiple regression. *Psychological Bulletin*, 114(3), 542–551. <https://doi.org/10.1037/0033-2909.114.3.542>
- Gu, X. (2022). Assessing the relative importance of predictors in latent regression models. *Structural Equation Modeling: A Multidisciplinary Journal*, 4, 569–583. <https://doi.org/10.1080/10705511.2021.2025377>

See Also

[dominance](#), [coeff.std](#), [write.result](#)

Examples

```

#-----
# Linear model

# Example 1a: Dominance analysis, 'mpg' predicted by 'cyl', 'disp', and 'hp'
dominance.manual(cor(mtcars[, c("mpg", "cyl", "disp", "hp")]))

# Example 1b: Equivalent results using the dominance() function
mod <- lm(mpg ~ cyl + disp + hp, data = mtcars)
dominance(mod)

# Example 1c: Dominance analysis, 'hp' predicted by 'mpg', 'cyl', and 'disp'
dominance.manual(cor(mtcars[, c("mpg", "cyl", "disp", "hp")]), out = "hp")

## Not run:
# Example 1d: Write results into a text file
dominance.manual(cor(mtcars[, c("mpg", "cyl", "disp", "hp")]),
  write = "Dominance_Manual.txt")
## End(Not run)

#-----
# Example 2: Structural equation modeling

library(lavaan)

#.....
# Latent variables

# Model specification
model <- '# Measurement model
  ind60 =~ x1 + x2 + x3
  dem60 =~ y1 + y2 + y3 + y4
  dem65 =~ y5 + y6 + y7 + y8
  # regressions
  ind60 ~ dem60 + dem65'

# Model estimation
fit <- sem(model, data = PoliticalDemocracy)

# Model-implied correlation matrix of the latent variables
fit.cor <- lavInspect(fit, what = "cor.lv")

# Dominance analysis
dominance.manual(fit.cor)

#.....
# Example 3: Latent and manifest variables

# Model specification, convert manifest to latent variable
model <- '# Measurement model
  ind60 =~ x1 + x2 + x3
  dem60 =~ y1 + y2 + y3 + y4

```

```

# Manifest as latent variable
ly5 =~ 1*y5
y5 ~~ 0*y5
# Regressions
ind60 ~ dem60 + ly5'

# Model estimation
fit <- sem(model, data = PoliticalDemocracy)

# Model-implied correlation matrix of the latent variables
fit.cor <- lavInspect(fit, what = "cor.lv")

# Dominance analysis
dominance.manual(fit.cor)

#-----
# Example 4: Multilevel modeling

# Model specification
model <- 'level: 1
    fw =~ y1 + y2 + y3
    # Manifest as latent variables
    lx1 =~ 1*x1
    lx2 =~ 1*x2
    lx3 =~ 1*x3
    x1 ~~ 0*x1
    x2 ~~ 0*x2
    x3 ~~ 0*x3
    # Regression
    fw ~ lx1 + lx2 + lx3
level: 2
    fb =~ y1 + y2 + y3
    # Manifest as latent variables
    lw1 =~ 1*w1
    lw2 =~ 1*w2
    # Regression
    fb ~ lw1 + lw2'

# Model estimation
fit <- sem(model, data = Demo.twolevel, cluster = "cluster")

# Model-implied correlation matrix of the latent variables
fit.cor <- lavInspect(fit, what = "cor.lv")

# Dominance analysis Within
dominance.manual(fit.cor$within)

# Dominance analysis Between
dominance.manual(fit.cor$cluster)

## Not run:
#-----
# Example 5: Mplus

```

```
#
# In Mplus, the model-implied correlation matrix of the latent variables
# can be requested by OUTPUT: TECH4 and imported into R by using the
# MplusAutomation package, for example:

library(MplusAutomation)

# Read Mplus output
output <- readModels()

# Extract model-implied correlation matrix of the latent variables
fit.cor <- output$tech4$latCorEst

## End(Not run)
```

effsize

Effect Sizes for Categorical Variables

Description

This function computes effect sizes for one or more than one categorical variable, i.e., (adjusted) phi coefficient, (bias-corrected) Cramer's V , (bias-corrected) Tschuprow's T , (adjusted) Pearson's contingency coefficient, Cohen's w , and *Fei*. By default, the function computes *Fei* based on a chi-square goodness-of-fit test for one categorical variable, phi coefficient based on a chi-square test of independence for two dichotomous variables, and Cramer's V based on a chi-square test of independence for two variables with at least one polytomous variable.

Usage

```
effsize(data, ..., type = c("phi", "cramer", "tschuprow", "cont", "w", "fei"),
        alternative = c("two.sided", "less", "greater"), conf.level = 0.95,
        adjust = TRUE, indep = TRUE, p = NULL, digits = 3, as.na = NULL,
        write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a vector, factor or data frame.
<code>...</code>	an expression indicating the variable names in <code>data</code> , e.g., <code>effsize(dat, x1, x2)</code> . When specifying more than one variable, the first variable is always the focal variable in the Chi-square test of independence which association with all other variables is investigated. Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>type</code>	a character string indicating the type of effect size, i.e., <code>phi</code> for phi coefficient, <code>cramer</code> for Cramer's V , <code>tschuprow</code> for Tschuprow's T , <code>cont</code> for Pearson's contingency coefficient, <code>w</code> for Cohen's w , and <code>Fei</code> for <i>Fei</i> .
<code>alternative</code>	a character string specifying the alternative hypothesis, must be one of <code>"two.sided"</code> (default), <code>"greater"</code> or <code>"less"</code> .

<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.
<code>adjust</code>	logical: if TRUE (default), phi coefficient and Pearson's contingency coefficient are adjusted by relating the coefficient to the possible maximum, or Cramer's V and Tschuprow's T are corrected for small-sample bias.
<code>indep</code>	logical: if TRUE, effect size computation is based on a chi-square test of independence (default when specifying two variable, if FALSE effect size computation is based on a chi-square goodness-of-fit test (default when specifying one variable).
<code>p</code>	a numeric vector specifying the expected proportions in each category of the categorical variable when conducting a chi-square goodness-of-fit test. By default, the expected proportions in each category are assumed to be equal.
<code>digits</code>	an integer value indicating the number of decimal places digits to be used for displaying the results.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame with variables used in the current analysis
<code>args</code>	specification of function arguments
<code>result</code>	result table

Note

This function is based on modified copies of the functions `chisq_to_phi`, `chisq_to_cramers_v`, `chisq_to_tschuprows_t`, `chisq_to_pearsons_c`, `chisq_to_cohens_w`, and `chisq_to_fei` from the **effectsize** package (Ben-Shachar, Lüdtke & Makowski, 2020).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Bergsma, W. (2013). A bias correction for Cramer's V and Tschuprow's T. *Journal of the Korean Statistical Society*, 42, 323-328. <https://doi.org/10.1016/j.jkss.2012.10.002>
- Ben-Shachar M. S., Lüdtke D., Makowski D. (2020). effectsize: Estimation of Effect Size Indices and Standardized Parameters. *Journal of Open Source Software*, 5 (56), 2815. <https://doi.org/10.21105/joss.02815>
- Ben-Shachar, M. S., Patil, I., Theriault, R., Wiernik, B. M., Lüdtke, D. (2023). Phi, Fei, Fo, Fum: Effect sizes for categorical data that use the chi-squared statistic. *Mathematics*, 11, 1982. <https://doi.org/10.3390/math11091982>
- Cureton, E. E. (1959). Note on Phi/Phi max. *Psychometrika*, 24, 89-91.
- Davenport, E. C., & El-Sanhuri, N. A. (1991). Phi/Phimax: Review and synthesis. *Educational and Psychological Measurement*, 51, 821-828. <https://doi.org/10.1177/001316449105100403>
- Sakoda, J.M. (1977). Measures of association for multivariate contingency tables. *Proceedings of the Social Statistics Section of the American Statistical Association (Part III)*, 777-780.

See Also

[cor.matrix](#), [cohens.d](#)

Examples

```
# Example 1: Phi coefficient for 'vs' and 'am'
effsize(mtcars, vs, am)

# Alternative specification without using the '...' argument
effsize(mtcars[, c("vs", "am")])

# Example 2: Bias-corrected Cramer's V for 'gear' and 'carb'
effsize(mtcars, gear, carb)

# Example 3: Cramer's V (without bias-correction) for 'gear' and 'carb'
effsize(mtcars, gear, carb, adjust = FALSE)

# Example 4: Adjusted Pearson's contingency coefficient for 'gear' and 'carb'
effsize(mtcars, gear, carb, type = "cont")

# Example 5: Fei for 'gear'
effsize(mtcars, gear)

# Example 6: Bias-corrected Cramer's V for 'cyl' and 'vs', 'am', 'gear', and 'carb'
effsize(mtcars, cyl, vs:carb)

# Alternative specification without using the '...' argument
effsize(mtcars[, c("cyl", "vs", "am", "gear", "carb")])

## Not run:
# Example 7a: Write Results into a text file
effsize(mtcars, cyl, vs:carb, write = "Cramer.txt")

# Example 7b: Write Results into a Excel file
```

```
effsize(mtcars, cyl, vs:carb, write = "Cramer.xlsx")

## End(Not run)
```

freq	<i>Frequency Table</i>
------	------------------------

Description

This function computes a frequency table with absolute and percentage frequencies for one or more than one variable. By default, the function displays the absolute and percentage frequencies when specifying one variable, while the function displays only the absolute frequencies when specifying more than one variable.

Usage

```
freq(data, ..., print = c("no", "all", "perc", "v.perc"), freq = TRUE,
      split = FALSE, labels = TRUE, val.col = FALSE, round = 3, exclude = 15,
      digits = 2, as.na = NULL, write = NULL, append = TRUE, check = TRUE,
      output = TRUE)
```

Arguments

data	a vector, factor, or data frame.
...	an expression indicating the variable names in data, e.g., <code>freq(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the <code>df.subset</code> function.
print	a character string indicating which percentage(s) to be printed on the console, i.e., no percentages ("no"), all percentages ("all"), percentage frequencies ("print"), and valid percentage frequencies ("v.perc"). Default setting when specifying one variable is <code>print = "all"</code> , while default setting when specifying more than one variable is <code>print = "no"</code> unless <code>split = TRUE</code> .
freq	logical: if TRUE (default), absolute frequencies will be shown on the console.
split	logical: if TRUE, output table is split by variables when specifying more than one variable in
labels	logical: if TRUE (default), labels for the factor levels will be used.
val.col	logical: if TRUE, values are shown in the columns, variables in the rows.
round	an integer value indicating the number of decimal places to be used for rounding numeric variables.
exclude	an integer value indicating the maximum number of unique values for variables to be included in the analysis when specifying more than one variable i.e., variables with the number of unique values exceeding <code>exclude</code> will be excluded from the analysis. It is also possible to specify <code>exclude = FALSE</code> to include all variables in the analysis.

<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying percentages.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

The function displays valid percentage frequencies only in the presence of missing values and excludes variables with all values missing from the analysis. Note that it is possible to mix numeric variables, factors, and character variables in the data frame specified in the argument `data`. By default, numeric variables are rounded to three digits before computing the frequency table.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame used for the current analysis
<code>args</code>	specification of function arguments
<code>result</code>	data frame with absolute frequencies and percentages or list with result tables, i.e., <code>freq</code> for absolute frequencies, <code>perc</code> for percentages, and <code>v.perc</code> for valid percentages

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Becker, R. A., Chambers, J. M., & Wilks, A. R. (1988). *The New S Language*. Wadsworth & Brooks/Cole.

See Also

[write.result](#), [crosstab](#), [descript](#), [multilevel.descript](#), [na.descript](#).

Examples

```
# Example 1: Frequency table for 'cyl'
freq(mtcars, cyl)

# Alternative specification without using the '...' argument
freq(mtcars$cyl)

# Example 2: Frequency table, values shown in columns
freq(mtcars, cyl, val.col = TRUE)

# Example 3: Frequency table, use 3 digit for displaying percentages
freq(mtcars, cyl, digits = 3)

# Example 4: Frequency table for 'cyl', 'gear', and 'carb'
freq(mtcars, cyl, gear, carb)

# Alternative specification without using the '...' argument
freq(mtcars[, c("cyl", "gear", "carb")])

# Example 5: Frequency table, with percentage frequencies
freq(mtcars, cyl, gear, carb, print = "all")

# Example 6: Frequency table, split output table
freq(mtcars, cyl, gear, carb, split = TRUE)

# Example 7: Frequency table, exclude variables with more than 5 unique values
freq(mtcars, exclude = 5)

## Not run:
# Example 8a: Write Results into a text file
freq(mtcars, cyl, gear, carb, split = TRUE, write = "Frequencies.txt")

# Example 8b: Write Results into a Excel file
freq(mtcars, cyl, gear, carb, split = TRUE, write = "Frequencies.xlsx")
## End(Not run)
```

indirect

Confidence Intervals for the Indirect Effect

Description

This function computes confidence intervals for the indirect effect based on the asymptotic normal method, distribution of the product method and the Monte Carlo method. By default, the function uses the distribution of the product method for computing the two-sided 95% asymmetric confidence intervals for the indirect effect product of coefficient estimator $\hat{a}\hat{b}$.

Usage

```
indirect(a, b, se.a, se.b, print = c("all", "asyp", "dop", "mc"),
```

```
se = c("sobel", "aroian", "goodman"), nrep = 100000,
alternative = c("two.sided", "less", "greater"), seed = NULL,
conf.level = 0.95, digits = 3, write = NULL, append = TRUE,
check = TRUE, output = TRUE)
```

Arguments

<code>a</code>	a numeric value indicating the coefficient a , i.e., effect of X on M .
<code>b</code>	a numeric value indicating the coefficient b , i.e., effect of M on Y adjusted for X .
<code>se.a</code>	a positive numeric value indicating the standard error of a .
<code>se.b</code>	a positive numeric value indicating the standard error of b .
<code>print</code>	a character string or character vector indicating which confidence intervals (CI) to show on the console, i.e. "all" for all CIs, "asyp" for the CI based on the asymptotic normal method, "dop" (default) for the CI based on the distribution of the product method, and "mc" for the CI based on the Monte Carlo method.
<code>se</code>	a character string indicating which standard error (SE) to compute for the asymptotic normal method, i.e., "sobel" for the approximate standard error by Sobel (1982) using the multivariate delta method based on a first order Taylor series approximation, "aroian" (default) for the exact standard error by Aroian (1947) based on a first and second order Taylor series approximation, and "goodman" for the unbiased standard error by Goodman (1960).
<code>nrep</code>	an integer value indicating the number of Monte Carlo repetitions.
<code>alternative</code>	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
<code>seed</code>	a numeric value specifying the seed of the random number generator when using the Monte Carlo method.
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying
<code>write</code>	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

In statistical mediation analysis (MacKinnon & Tofighi, 2013), the indirect effect refers to the effect of the independent variable X on the outcome variable Y transmitted by the mediator variable M . The magnitude of the indirect effect ab is quantified by the product of the coefficient a (i.e., effect of X on M) and the coefficient b (i.e., effect of M on Y adjusted for X). In practice, researchers are often interested in confidence limit estimation for the indirect effect. This function

offers three different methods for computing the confidence interval for the product of coefficient estimator $\hat{a}\hat{b}$:

(1) Asymptotic normal method

In the asymptotic normal method, the standard error for the product of the coefficient estimator $\hat{a}\hat{b}$ is computed which is used to create a symmetrical confidence interval based on the z-value of the standard normal (z) distribution assuming that the indirect effect is normally distributed. Note that the function provides three formulas for computing the standard error by specifying the argument `se`:

"sobel" Approximate standard error by Sobel (1982) using the multivariate delta method based on a first order Taylor series approximation:

$$\sqrt{(a^2\sigma_a^2 + b^2\sigma_b^2)}$$

"aroian" Exact standard error by Aroian (1947) based on a first and second order Taylor series approximation:

$$\sqrt{(a^2\sigma_a^2 + b^2\sigma_b^2 + \sigma_a^2\sigma_b^2)}$$

"goodman" Unbiased standard error by Goodman (1960):

$$\sqrt{(a^2\sigma_a^2 + b^2\sigma_b^2 - \sigma_a^2\sigma_b^2)}$$

Note that the unbiased standard error is often negative and is hence undefined for zero or small effects or small sample sizes.

The asymptotic normal method is known to have low statistical power because the distribution of the product $\hat{a}\hat{b}$ is not normally distributed. (Kisbu-Sakarya, MacKinnon, & Miocevic, 2014). In the null case, where both random variables have mean equal to zero, the distribution is symmetric with kurtosis of six. When the product of the means of the two random variables is nonzero, the distribution is skewed (up to a maximum value of ± 1.5) and has a excess kurtosis (up to a maximum value of 6). However, the product approaches a normal distribution as one or both of the ratios of the means to standard errors of each random variable get large in absolute value (MacKinnon, Lockwood & Williams, 2004).

(2) Distribution of the product method

The distribution of the product method (MacKinnon et al., 2002) relies on an analytical approximation of the distribution of the product of two normally distributed variables. The method uses the standardized a and b coefficients to compute ab and then uses the critical values for the distribution of the product (Meeker, Cornwell, & Aroian, 1981) to create asymmetric confidence intervals. The distribution of the product approaches the gamma distribution (Aroian, 1947). The analytical solution for the distribution of the product is provided by the Bessel function used to the solution of differential equations and is approximately proportional to the Bessel function of the second kind with a purely imaginary argument (Craig, 1936).

(3) Monte Carlo method

The Monte Carlo (MC) method (MacKinnon et al., 2004) relies on the assumption that the parameters a and b have a joint normal sampling distribution. Based on the parametric assumption, a sampling distribution of the product ab using random samples with population values equal to the sample estimates $\hat{a}$, $\hat{b}$, $\hat{\sigma}_a$, and $\hat{\sigma}_b$ is generated. Percentiles of the sampling distribution are identified to serve as limits for a $100(1 - \alpha)\%$ asymmetric confidence interval about the sample $\hat{a}\hat{b}$ (Preacher & Selig, 2012). Note that parametric assumptions are invoked for $\hat{a}$ and $\hat{b}$, but no parametric assumptions are made about the distribution of $\hat{a}\hat{b}$.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	list with the input specified in <code>a.b</code> , <code>se.a</code> , and <code>se.b</code>
<code>args</code>	specification of function arguments
<code>result</code>	list with result tables, i.e., <code>asyp</code> with CI based on the asymptotic normal method, <code>dop</code> with CI based on the distribution of the product method, and <code>mc</code> for CI based on the Monte Carlo method

Note

The function was adapted from the `medci()` function in the **RMediation** package by Davood Tofighi and David P. MacKinnon (2016).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Aroian, L. A. (1947). The probability function of the product of two normally distributed variables. *Annals of Mathematical Statistics*, 18, 265-271. <https://doi.org/10.1214/aoms/1177730442>
- Craig, C.C. (1936). On the frequency function of xy . *Annals of Mathematical Statistics*, 7, 1-15. <https://doi.org/10.1214/aoms/1177732541>
- Goodman, L. A. (1960). On the exact variance of products. *Journal of the American Statistical Association*, 55, 708-713. <https://doi.org/10.1080/01621459.1960.10483369>
- Kisbu-Sakarya, Y., MacKinnon, D. P., & Miocevic M. (2014). The distribution of the product explains normal theory mediation confidence interval estimation. *Multivariate Behavioral Research*, 49, 261-268. <https://doi.org/10.1080/00273171.2014.903162>
- MacKinnon, D. P., Lockwood, C. M., Hoffman, J. M., West, S. G., & Sheets, V. (2002). Comparison of methods to test mediation and other intervening variable effects. *Psychological Methods*, 7, 83-104. <https://doi.org/10.1037/1082-989x.7.1.83>
- MacKinnon, D. P., Lockwood, C. M., & Williams, J. (2004). Confidence limits for the indirect effect: Distribution of the product and resampling methods. *Multivariate Behavioral Research*, 39, 99-128. https://doi.org/10.1207/s15327906mbr3901_4
- MacKinnon, D. P., & Tofighi, D. (2013). Statistical mediation analysis. In J. A. Schinka, W. F. Velicer, & I. B. Weiner (Eds.), *Handbook of psychology: Research methods in psychology* (pp. 717-735). John Wiley & Sons, Inc..
- Meeker, W. Q., Jr., Cornwell, L. W., & Aroian, L. A. (1981). The product of two normally distributed random variables. In W. J. Kennedy & R. E. Odeh (Eds.), *Selected tables in mathematical statistics* (Vol. 7, pp. 1-256). Providence, RI: American Mathematical Society.
- Preacher, K. J., & Selig, J. P. (2012). Advantages of Monte Carlo confidence intervals for indirect effects. *Communication Methods and Measures*, 6, 77-98. <http://dx.doi.org/10.1080/19312458.2012.679848>

Sobel, M. E. (1982). Asymptotic confidence intervals for indirect effects in structural equation models. In S. Leinhardt (Ed.), *Sociological methodology 1982* (pp. 290-312). Washington, DC: American Sociological Association.

Tofighi, D. & MacKinnon, D. P. (2011). RMediation: An R package for mediation analysis confidence intervals. *Behavior Research Methods*, 43, 692-700. <https://doi.org/10.3758/s13428-011-0076-x>

See Also

[multilevel.indirect](#)

Examples

```
# Example 1: Distribution of the Product Method
indirect(a = 0.35, b = 0.27, se.a = 0.12, se.b = 0.18)

# Example 2: Monte Carlo Method
indirect(a = 0.35, b = 0.27, se.a = 0.12, se.b = 0.18, print = "mc")

# Example 3: Asymptotic Normal Method
indirect(a = 0.35, b = 0.27, se.a = 0.12, se.b = 0.18, print = "asymp")

## Not run:
# Example 4: Write results into a text file
indirect(a = 0.35, b = 0.27, se.a = 0.12, se.b = 0.18, write = "Indirect.txt")
## End(Not run)
```

item.alpha

Coefficient Alpha, Hierarchical Alpha, and Ordinal Alpha

Description

This function computes point estimate and confidence interval for the coefficient alpha (aka Cronbach's alpha), hierarchical alpha, and ordinal alpha (aka categorical alpha) along with standardized factor loadings and alpha if item deleted. By default, the function computes coefficient alpha based on unweighted least squares (ULS) parameter estimates using pairwise deletion in the presence of missing data that provides equivalent results compared to the formula-based coefficient alpha computed by using e.g. the alpha function in the **psych** package by William Revelle (2025).

Usage

```
item.alpha(data, ..., rescov = NULL, type = c("alpha", "hierarch", "categ"),
  exclude = NULL, std = FALSE,
  estimator = c("ML", "GLS", "WLS", "DWLS", "ULS", "PML"),
  missing = c("listwise", "pairwise", "fiml"),
  print = c("all", "alpha", "item"), digits = 2, conf.level = 0.95,
  as.na = NULL, write = NULL, append = TRUE, check = TRUE,
  output = TRUE)
```

Arguments

data	a data frame. Note that at least two items are needed for computing coefficient alpha
...	an expression indicating the variable names in data e.g., <code>item.alpha(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the <code>df.subset</code> function.
rescov	a character vector or a list of character vectors for specifying residual covariances when computing coefficient alpha, e.g. <code>rescov = c("x1", "x2")</code> for specifying a residual covariance between items x1 and x2 or <code>rescov = list(c("x1", "x2"), c("x3", "x4"))</code> for specifying residual covariances between items x1 and x2, and items x3 and x4.
type	a character string indicating the type of alpha to be computed, i.e., <code>alpha</code> (default) for coefficient alpha, <code>hierarch</code> for hierarchical coefficient alpha, and <code>categ</code> for ordinal coefficient alpha.
exclude	a character vector indicating items to be excluded from the analysis.
std	logical: if TRUE, the standardized coefficient omega is computed.
estimator	a character string indicating the estimator to be used (see 'Details' in the <code>item.cfa</code> function). By default, "ULS" is used for computing (hierarchical) coefficient alpha and "DWLS" is used for computing ordinal coefficient alpha.
missing	a character string indicating how to deal with missing data. (see 'Details' in the <code>item.cfa</code> function). By default, pairwise deletion (<code>missing = "pairwise"</code>) is used for computing (hierarchical) coefficient alpha and ordinal coefficient alpha. Full information maximum likelihood method is available for estimating (hierarchical) coefficient alpha and is requested by specifying <code>missing = "fiml"</code> along with <code>estimator = "ML"</code> .
print	a character vector indicating which results to show, i.e. <code>"all"</code> for all results <code>"alpha"</code> (default) for the coefficient alpha, and <code>"item"</code> for item statistics.
digits	an integer value indicating the number of decimal places to be used for displaying alpha and standardized factor loadings.
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
write	a character string naming a file for writing the output into either a text file with file extension <code>".txt"</code> (e.g., <code>"Output.txt"</code>) or Excel file with file extension <code>".xlsx"</code> (e.g., <code>"Output.xlsx"</code>). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension <code>.txt</code> specified in <code>write</code> , if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown.

Details

Coefficient alpha is computed by conducting a confirmatory factor analysis based on the essentially tau-equivalent measurement model (Graham, 2006) using the `cfa()` function in the **lavaan** package by Yves Rosseel (2019). Approximate confidence intervals are computed using the procedure by Feldt, Woodruff and Salih (1987). Note that there are at least 10 other procedures for computing the confidence interval (see Kelley and Pornprasertmanit, 2016), which are implemented in the `ci.reliability()` function in the **MBESS** package by Ken Kelley (2019)

Ordinal coefficient alpha was introduced by Zumbo, Gadermann and Zeisser (2007). Note that Chalmers (2018) highlighted that the categorical coefficient alpha should be interpreted only as a hypothetical estimate of an alternative reliability, whereby a test's ordinal categorical response options have been modified to include an infinite number of ordinal response options and concludes that coefficient alpha should not be reported as a measure of a test's reliability. However, Zumbo and Kroc (2019) argued that Chalmers' critique of categorical coefficient alpha is unfounded and that categorical coefficient alpha may be the most appropriate quantifier of reliability when using Likert-type measurement to study a latent continuous random variable.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame used for the current analysis
<code>args</code>	specification of function arguments
<code>model.fit</code>	fitted lavaan object (<code>mod.fit</code>)
<code>result</code>	list with result tables, i.e., alpha for a table with coefficient alpha and <code>itemstat</code> for a table with item statistics

Note

Computation of the hierarchical and ordinal alpha is based on the `ci.reliability()` function in the **MBESS** package by Ken Kelley (2019).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Chalmers, R. P. (2018). On misconceptions and the limited usefulness of ordinal alpha. *Educational and Psychological Measurement*, 78, 1056-1071. <https://doi.org/10.1177/0013164417727036>
- Cronbach, L.J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16, 297-334. <https://doi.org/10.1007/BF02310555>
- Cronbach, L.J. (2004). My current thoughts on coefficient alpha and successor procedures. *Educational and Psychological Measurement*, 64, 391-418. <https://doi.org/10.1177/0013164404266386>
- Feldt, L. S., Woodruff, D. J., & Salih, F. A. (1987). Statistical inference for coefficient alpha. *Applied Psychological Measurement*, 11 93-103. <https://doi.org/10.1177/014662168701100107>

Graham, J. M. (2006). Congeneric and (essentially) tau-equivalent estimates of score reliability: What they are and how to use them. *Educational and Psychological Measurement*, 66(6), 930–944. <https://doi.org/10.1177/0013164406288165>

Kelley, K., & Pornprasertmanit, S. (2016). Confidence intervals for population reliability coefficients: Evaluation of methods, recommendations, and software for composite measures. *Psychological Methods*, 21, 69–92. <https://doi.org/10.1037/a0040086>.

Ken Kelley (2019). *MBESS: The MBESS R Package*. R package version 4.6.0. <https://CRAN.R-project.org/package=MBESS>

Revelle, W. (2025). *psych: Procedures for psychological, psychometric, and personality research*. Northwestern University, Evanston, Illinois. R package version 2.5.3, <https://CRAN.R-project.org/package=psych>.

Zumbo, B. D., & Kroc, E. (2019). A measurement is a choice and Stevens' scales of measurement do not help make it: A response to Chalmers. *Educational and Psychological Measurement*, 79, 1184–1197. <https://doi.org/10.1177/0013164419844305>

Zumbo, B. D., Gadermann, A. M., & Zeisser, C. (2007). Ordinal versions of coefficients alpha and theta for Likert rating scales. *Journal of Modern Applied Statistical Methods*, 6, 21–29. <https://doi.org/10.22237/jmasm/1177992180>

See Also

[item.omega](#), [item.cfa](#), [item.invar](#), [item.reverse](#), [item.scores](#), [write.result](#)

Examples

```
## Not run:
dat <- data.frame(item1 = c(3, NA, 3, 4, 1, 2, 4, 2), item2 = c(5, 3, 3, 2, 2, 1, 3, 1),
                  item3 = c(4, 2, 4, 2, 1, 3, 4, 1), item4 = c(4, 1, 2, 2, 1, 3, 4, 3))

# Example 1a: Coefficient alpha, pairwise deletion
item.alpha(dat)

# Example 1b: Coefficient alpha, listwise deletion
item.alpha(dat, missing = "listwise")

# Example 1c: Coefficient alpha, Full information maximum likelihood method
item.alpha(dat, estimator = "ML", missing = "fiml")

# Example 2: Coefficient alpha and item statistics after excluding item3
item.alpha(dat, exclude = "item3", print = "all")

# Example 3a: Coefficient alpha with a residual covariance
item.alpha(dat, rescov = c("item1", "item2"))

# Example 3b: Coefficient alpha with residual covariances
item.alpha(dat, rescov = list(c("item1", "item2"), c("item1", "item3")))

# Example 4: Ordinal coefficient alpha and item statistics
item.alpha(dat, type = "categ", print = "all")

# Example 6: Summary of the CFA model used to compute coefficient alpha
lavaan::summary(item.alpha(dat, output = FALSE)$model.fit,
```

```

fit.measures = TRUE, standardized = TRUE)

# Example 7a: Write Results into a text file
item.alpha(dat, write = "Alpha.txt")

# Example 7b: Write Results into a Excel file
item.alpha(dat, write = "Alpha.xlsx")

## End(Not run)

```

item.cfa

Confirmatory Factor Analysis

Description

This function is a wrapper function for conducting confirmatory factor analysis with continuous and/or ordered-categorical indicators by calling the `cfa` function in the R package **lavaan**.

Usage

```

item.cfa(data, ..., model = NULL, rescov = NULL, hierarch = FALSE,
  meanstructure = TRUE, ident = c("marker", "var", "effect"),
  parameterization = c("delta", "theta"), ordered = NULL, cluster = NULL,
  estimator = c("ML", "MLM", "MLMV", "MLMVS", "MLF", "MLR",
    "GLS", "WLS", "DWLS", "WLSM", "WLSMV",
    "ULS", "ULSM", "ULSMV", "DLS", "PML"),
  missing = c("listwise", "pairwise", "fiml",
    "two.stage", "robust.two.stage", "doubly.robust"),
  print = c("all", "summary", "coverage", "descript", "fit", "est",
    "modind", "resid"),
  mod.minval = 6.63, resid.minval = 0.1, digits = 3, p.digits = 3,
  as.na = NULL, write = NULL, append = TRUE, check = TRUE, output = TRUE)

```

Arguments

<code>data</code>	a data frame. If <code>model = NULL</code> , confirmatory factor analysis based on a measurement model with one factor labeled <code>f</code> comprising all variables in the data frame is conducted. Note that the cluster variable is excluded from data when specifying cluster. If <code>model</code> is specified, the data frame needs to contain all variables used in the argument <code>model</code> and the cluster variable when specifying cluster.
<code>...</code>	an expression indicating the variable names in <code>data</code> , e.g., <code>item.cfa(x1, x2, x3, data = dat)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>model</code>	a character vector specifying a measurement model with one factor, or a list of character vectors for specifying a measurement model with more than one factor, e.g., <code>model = c("x1", "x2", "x3", "x4")</code> for specifying a measurement model

	with one factor labeled <code>f</code> comprising four indicators, or <code>model = list(factor1 = c("x1", "x2", "x3", "x4"), factor2 = c("x5", "x6", "x7", "x8"))</code> for specifying a measurement model with two latent factors labeled <code>factor1</code> and <code>factor2</code> each comprising four indicators. Note that the name of each list element is used to label factors, i.e., all list elements need to be named, otherwise factors are labeled with <code>"f1"</code> , <code>"f2"</code> , <code>"f3"</code> and so on.
<code>rescov</code>	a character vector or a list of character vectors for specifying residual covariances, e.g. <code>rescov = c("x1", "x2")</code> for specifying a residual covariance between items <code>x1</code> and <code>x2</code> , or <code>rescov = list(c("x1", "x2"), c("x3", "x4"))</code> for specifying residual covariances between items <code>x1</code> and <code>x2</code> , and items <code>x3</code> and <code>x4</code> .
<code>hierarchy</code>	logical: if <code>TRUE</code> , a second-order factor model is specified given at least three first-order factors were specified in <code>model</code> . Note that it is not possible to specify more than one second-order factor.
<code>meanstructure</code>	logical: if <code>TRUE</code> (default), intercept/means of observed variables means of latent variables will be added to the model. Note that <code>meanstructure = FALSE</code> is only applicable when the missing is <code>listwise</code> , <code>pairwise</code> , or <code>doubly-robust</code> .
<code>ident</code>	a character string indicating the method used for identifying and scaling latent variables, i.e., <code>"marker"</code> for the marker variable method fixing the first factor loading of each latent variable to 1, <code>"var"</code> for the fixed variance method fixing the variance of each latent variable to 1, or <code>"effect"</code> for the effects-coding method using equality constraints so that the average of the factor loading for each latent variable equals 1. By default, fixed variance method is used when <code>hierarchy = FALSE</code> , whereas marker variable method is used when <code>hierarchy = TRUE</code> .
<code>parameterization</code>	a character string indicating the method used for identifying and scaling latent variables when indicators are ordered, i.e., <code>"delta"</code> (default) for delta parameterization and <code>"theta"</code> for theta parameterization.
<code>ordered</code>	if <code>NULL</code> (default), all indicators of the measurement model are treated as continuous. If <code>TRUE</code> , all indicators of the measurement model are treated as ordered (ordinal). Alternatively, a character vector indicating which variables to treat as ordered (ordinal) variables can be specified.
<code>cluster</code>	either a character string indicating the variable name of the cluster variable in data, or a vector representing the nested grouping structure (i.e., group or cluster variable) for computing cluster-robust standard errors. Note that cluster-robust standard errors are not available when treating indicators of the measurement model as ordered (ordinal).
<code>estimator</code>	a character string indicating the estimator to be used (see 'Details'). By default, <code>"MLR"</code> is used for CFA models with continuous indicators (i.e., <code>ordered = FALSE</code>) and <code>"WLSMV"</code> is used for CFA model with ordered-categorical indicators (i.e., <code>ordered = TRUE</code>).
<code>missing</code>	a character string indicating how to deal with missing data, i.e., <code>"listwise"</code> for listwise deletion, <code>"pairwise"</code> for pairwise deletion, <code>"fiml"</code> for full information maximum likelihood method, <code>two.stage</code> for two-stage maximum likelihood method, <code>robust.two.stage</code> for robust two-stage maximum likelihood method, and <code>doubly-robust</code> for doubly-robust method (see 'Details').

	By default, "fiml" is used for CFA models with continuous indicators which are estimated by using <code>estimator = "MLR"</code> , and "pairwise" for CFA models with ordered-categorical indicators which are estimated by using <code>estimator = "pairwise"</code> by default.
<code>print</code>	a character string or character vector indicating which results to show on the console, i.e. "all" for all results, "summary" for a summary of the specification of the estimation method and missing data handling in lavaan, "coverage" for the variance-covariance coverage of the data, "descript" for descriptive statistics, "fit" for model fit, "est" for parameter estimates, "modind" for modification indices and "resid" for the residual correlation matrix and standardized residual means. By default, a summary of the specification, model fit, and parameter estimates are printed. By default, a summary of the specification, model fit, and parameter estimates are printed.
<code>mod.minval</code>	numeric value to filter modification indices and only show modifications with a modification index value equal or higher than this minimum value. By default, modification indices equal or higher 6.63 are printed. Note that a modification index value of 6.63 is equivalent to a significance level of $\alpha = .01$.
<code>resid.minval</code>	numeric value indicating the minimum absolute residual correlation coefficients and standardized means to highlight in boldface. By default, absolute residual correlation coefficients and standardized means equal or higher 0.1 are highlighted. Note that highlighting can be disabled by setting the minimum value to 1.
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying results.
<code>p.digits</code>	an integer value indicating the number of decimal places to be used for displaying the <i>p</i> -value.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to data but not to cluster.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown.

Details

Estimator The R package **lavaan** provides seven estimators that affect the estimation, namely "ML", "GLS", "WLS", "DWLS", "ULS", "DLS", and "PML". All other options for the argument `estimator` combine these estimators with various standard error and chi-square test statistic computation. Note that the estimators also differ in how missing values can be dealt with (e.g., listwise deletion, pairwise deletion, or full information maximum likelihood, FIML).

- "ML": Maximum likelihood parameter estimates with conventional standard errors and conventional test statistic. For both complete and incomplete data using pairwise deletion or FIML.
- "MLM": Maximum likelihood parameter estimates with conventional robust standard errors and a Satorra-Bentler scaled test statistic that are robust to non-normality. For complete data only.
- "MLMV": Maximum likelihood parameter estimates with conventional robust standard errors and a mean and a variance adjusted test statistic using a scale-shifted approach that are robust to non-normality. For complete data only.
- "MLMVS": Maximum likelihood parameter estimates with conventional robust standard errors and a mean and a variance adjusted test statistic using the Satterthwaite approach that are robust to non-normality. For complete data only.
- "MLF": Maximum likelihood parameter estimates with standard errors approximated by first-order derivatives and conventional test statistic. For both complete and incomplete data using pairwise deletion or FIML.
- "MLR": Maximum likelihood parameter estimates with Huber-White robust standard errors a test statistic which is asymptotically equivalent to the Yuan-Bentler T2* test statistic that are robust to non-normality and non-independence of observed when specifying a cluster variable using the argument cluster. For both complete and incomplete data using pairwise deletion or FIML.
- "GLS": Generalized least squares parameter estimates with conventional standard errors and conventional test statistic that uses a normal-theory based weight matrix. For complete data only. and conventional chi-square test. For both complete and incomplete data.
- "WLS": Weighted least squares parameter estimates (sometimes called ADF estimation) with conventional standard errors and conventional test statistic that uses a full weight matrix. For complete data only.
- "DWLS": Diagonally weighted least squares parameter estimates which uses the diagonal of the weight matrix for estimation with conventional standard errors and conventional test statistic. For both complete and incomplete data using pairwise deletion.
- "WLSM": Diagonally weighted least squares parameter estimates which uses the diagonal of the weight matrix for estimation, but uses the full weight matrix for computing the conventional robust standard errors and a Satorra-Bentler scaled test statistic. For both complete and incomplete data using pairwise deletion.
- "WLSMV": Diagonally weighted least squares parameter estimates which uses the diagonal of the weight matrix for estimation, but uses the full weight matrix for computing the conventional robust standard errors and a mean and a variance adjusted test statistic using a scale-shifted approach. For both complete and incomplete data using pairwise deletion.
- "ULS": Unweighted least squares parameter estimates with conventional standard errors and conventional test statistic. For both complete and incomplete data using pairwise deletion.
- "ULSM": Unweighted least squares parameter estimates with conventional robust standard errors and a Satorra-Bentler scaled test statistic. For both complete and incomplete data using pairwise deletion.
- "ULSMV": Unweighted least squares parameter estimates with conventional robust standard errors and a mean and a variance adjusted test statistic using a scale-shifted approach. For both complete and incomplete data using pairwise deletion.

- "DLS": Distributionally-weighted least squares parameter estimates with conventional robust standard errors and a Satorra-Bentler scaled test statistic. For complete data only.
- "PML": Pairwise maximum likelihood parameter estimates with Huber-White robust standard errors and a mean and a variance adjusted test statistic using the Satterthwaite approach. For both complete and incomplete data using pairwise deletion.

Missing Data The R package **lavaan** provides six methods for dealing with missing data:

- "listwise": Listwise deletion, i.e., all cases with missing values are removed from the data before conducting the analysis. This is only valid if the data are missing completely at random (MCAR).
- "pairwise": Pairwise deletion, i.e., each element of a variance-covariance matrix is computed using cases that have data needed for estimating that element. This is only valid if the data are missing completely at random (MCAR).
- "fiml": Full information maximum likelihood (FIML) method, i.e., likelihood is computed case by case using all available data from that case. FIML method is only applicable for following estimators: "ML", "MLF", and "MLR".
- "two.stage": Two-stage maximum likelihood estimation, i.e., sample statistics is estimated using EM algorithm in the first step. Then, these estimated sample statistics are used as input for a regular analysis. Standard errors and test statistics are adjusted correctly to reflect the two-step procedure. Two-stage method is only applicable for following estimators: "ML", "MLF", and "MLR".
- "robust.two.stage": Robust two-stage maximum likelihood estimation, i.e., two-stage maximum likelihood estimation with standard errors and a test statistic that are robust against non-normality. Robust two-stage method is only applicable for following estimators: "ML", "MLF", and "MLR".
- "doubly.robust": Doubly-robust method only applicable for pairwise maximum likelihood estimation (i.e., estimator = "PML").

Convergence and model identification checks In line with the R package **lavaan**, this functions provides several checks for model convergence and model identification:

- Degrees of freedom: An error message is printed if the number of degrees of freedom is negative, i.e., the model is not identified.
- Model convergence: An error message is printed if the optimizer has not converged, i.e., results are most likely unreliable.
- Standard errors: An error message is printed if the standard errors could not be computed, i.e., the model might not be identified.
- Variance-covariance matrix of the estimated parameters: A warning message is printed if the variance-covariance matrix of the estimated parameters is not positive definite, i.e., the smallest eigenvalue of the matrix is smaller than zero or very close to zero.
- Negative variances of observed variables: A warning message is printed if the estimated variances of the observed variables are negative.
- Variance-covariance matrix of observed variables: A warning message is printed if the estimated variance-covariance matrix of the observed variables is not positive definite, i.e., the smallest eigenvalue of the matrix is smaller than zero or very close to zero.
- Negative variances of latent variables: A warning message is printed if the estimated variances of the latent variables are negative.

- Variance-covariance matrix of latent variables: A warning message is printed if the estimated variance-covariance matrix of the latent variables is not positive definite, i.e., the smallest eigenvalue of the matrix is smaller than zero or very close to zero.

Note that unlike the R package **lavaan**, the `item.cfa` function does not provide any results when the degrees of freedom is negative, the model has not converged, or standard errors could not be computed.

Model Fit The `item.cfa` function provides the chi-square test, incremental fit indices (i.e., CFI and TLI), and absolute fit indices (i.e., RMSEA, and SRMR) to evaluate overall model fit. However, different versions of the CFI, TLI, and RMSEA are provided depending on the estimator. Unlike the R package **lavaan**, the different versions are labeled with Standard, Scaled, and Robust in the output:

- "Standard": CFI, TLI, and RMSEA without any non-normality corrections. These fit measures based on the normal theory maximum likelihood test statistic are sensitive to deviations from multivariate normality of endogenous variables. Simulation studies by Brosseau-Liard et al. (2012), and Brosseau-Liard and Savalei (2014) showed that the uncorrected fit indices are affected by non-normality, especially at small and medium sample sizes (e.g., $n < 500$).
- "Scaled": Population-corrected robust CFI, TLI, and RMSEA with ad hoc non-normality corrections that simply replace the maximum likelihood test statistic with a robust test statistic (e.g., mean-adjusted chi-square). These fit indices change the population value being estimated depending on the degree of non-normality present in the data. Brosseau-Liard et al. (2012) demonstrated that the ad hoc corrected RMSEA increasingly accepts poorly fitting models as non-normality in the data increases, while the effect of the ad hoc correction on the CFI and TLI is less predictable with non-normality making fit appear worse, better, or nearly unchanged (Brosseau-Liard & Savalei, 2014).
- "Robust": Sample-corrected robust CFI, TLI, and RMSEA with non-normality corrections based on formula provided by Li and Bentler (2006) and Brosseau-Liard and Savalei (2014). These fit indices do not change the population value being estimated and can be interpreted the same way as the uncorrected fit indices when the data would have been normal.

In conclusion, the use of sample-corrected fit indices (Robust) instead of population-corrected fit indices (Scaled) is recommended. Note that when sample size is very small (e.g., $n < 200$), non-normality correction does not appear to adjust fit indices sufficiently to counteract the effect of non-normality (Brosseau-Liard & Savalei, 2014).

Modification Indices and Residual Correlation Matrix The `item.cfa` function provides modification indices and the residual correlation matrix when requested by using the `print` argument. Modification indices (aka score tests) are univariate Lagrange Multipliers (LM) representing a chi-square statistic with a single degree of freedom. LM approximates the amount by which the chi-square test statistic would decrease if a fixed or constrained parameter is freely estimated (Kline, 2023). However, (standardized) expected parameter change (EPC) values should also be inspected since modification indices are sensitive to sample size. EPC values are an estimate of how much the parameter would be expected to change if it were freely estimated (Brown, 2023). The residual correlation matrix is computed by separately converting the sample covariance and model-implied covariance matrices to correlation matrices before calculation differences between observed and predicted covariances (i.e., `type = "cor.bollen"`). As a rule of thumb, absolute correlation residuals greater than .10 indicate possible evidence for poor local fit, whereas smaller correlation residuals than 0.05 indicate

negligible degree of model misfit (Maydeu-Olivares, 2017). There is no reliable connection between the size of diagnostic statistics (i.e., modification indices and residuals) and the type or amount of model misspecification since (1) diagnostic statistics are themselves affected by misspecification, (2) misspecification in one part of the model distorts estimates in other parts of the model (i.e., error propagation), and (3) equivalent models have identical residuals but contradict the pattern of causal effects (Kline, 2023). Note that according to Kline' (2023) "any report of the results without information about the residuals is deficient" (p. 172).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame specified in data
<code>args</code>	specification of function arguments
<code>model</code>	specified model
<code>model.fit</code>	fitted lavaan object (<code>mod.fit</code>)
<code>check</code>	results of the convergence and model identification check
<code>result</code>	list with result tables, i.e., summary for the specification of the estimation method and missing data handling in lavaan, "coverage" for the variance-covariance coverage of the data, "descript" for descriptive statistics, <code>itemfreq</code> for absolute frequencies (<code>freq</code>), percentages (<code>perc</code>), and (<code>v.perc</code>) valid percentages, "fit" for model fit, "param" for parameter estimates, and "modind" for modification indices.

Note

The function uses the functions `cfa`, `lavInspect`, `lavTech`, `modindices`, `parameterEstimates`, and `standardizedsolution` provided in the R package **lavaan** by Yves Rosseel (2012).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Brosseau-Liard, P. E., Savalei, V., & Li, L. (2012). An investigation of the sample performance of two nonnormality corrections for RMSEA, *Multivariate Behavioral Research*, 47, 904-930. <https://doi.org/10.1080/00273171.2014.933697>
- Brosseau-Liard, P. E., & Savalei, V. (2014) Adjusting incremental fit indices for nonnormality. *Multivariate Behavioral Research*, 49, 460-470. <https://doi.org/10.1080/00273171.2014.933697>
- Brown, T. A. (2023). Confirmatory factor analysis. In R. H. Hoyle (Ed.), *Handbook of structural equation modeling* (2nd ed.) (pp. 361–379). The Guilford Press.
- Kline, R. B. (2023). *Principles and practice of structural equation modeling* (5th ed.). Guilford Press.

Li, L., & Bentler, P. M. (2006). Robust statistical tests for evaluating the hypothesis of close fit of misspecified mean and covariance structural models. *UCLA Statistics Preprint #506*. University of California.

Maydeu-Olivares, A. (2017). Assessing the size of model misfit in structural equation models. *Psychometrika*, 82(3), 533–558. <https://doi.org/10.1007/s11336-016-9552-7>

Rosseel, Y. (2012). lavaan: An R Package for Structural Equation Modeling. *Journal of Statistical Software*, 48, 1-36. <https://doi.org/10.18637/jss.v048.i02>

See Also

[item.alpha](#), [item.omega](#), [item.scores](#)

Examples

```
# Load data set "HolzingerSwineford1939" in the lavaan package
data("HolzingerSwineford1939", package = "lavaan")

#-----
# Measurement model with one factor

# Example 1a: Specification using the argument '...'
item.cfa(HolzingerSwineford1939, x1:x3)

# Example 1b: Alternative specification without using the '...' argument
item.cfa(HolzingerSwineford1939[, c("x1", "x2", "x3")])

# Example 1c: Alternative specification using the argument 'model'
item.cfa(HolzingerSwineford1939, model = c("x1", "x2", "x3"))

# Example 1e: Alternative specification using the argument 'model'
item.cfa(HolzingerSwineford1939, model = list(visual = c("x1", "x2", "x3")))

#-----
# Measurement model with three factors

# Example 2: Specification using the argument 'model'
item.cfa(HolzingerSwineford1939,
  model = list(visual = c("x1", "x2", "x3"),
    textual = c("x4", "x5", "x6"),
    speed = c("x7", "x8", "x9")))

#-----
# Residual covariances

# Example 3a: One residual covariance
item.cfa(HolzingerSwineford1939,
  model = list(visual = c("x1", "x2", "x3"),
    textual = c("x4", "x5", "x6"),
    speed = c("x7", "x8", "x9")),
  rescov = c("x1", "x2"))
```

```

# Example 3b: Two residual covariances
item.cfa(HolzingerSwineford1939,
  model = list(visual = c("x1", "x2", "x3"),
    textual = c("x4", "x5", "x6"),
    speed = c("x7", "x8", "x9")),
  rescov = list(c("x1", "x2"), c("x4", "x5")))

#-----
# Second-order factor model based on three first-order factors

# Example 4
item.cfa(HolzingerSwineford1939,
  model = list(visual = c("x1", "x2", "x3"),
    textual = c("x4", "x5", "x6"),
    speed = c("x7", "x8", "x9")), hierarch = TRUE)

#-----
# Measurement model with ordered-categorical indicators

# Example 5
item.cfa(round(HolzingerSwineford1939[, c("x4", "x5", "x6")]), ordered = TRUE)

#-----
# Cluster-robust standard errors

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

# Example 6a: Specification using the '...' argument
item.cfa(y4:y6, data = Demo.twolevel, cluster = "cluster")

# Example 6b: Alternative specification without using the '...' argument
item.cfa(Demo.twolevel[, c("y4", "y5", "y6")], cluster = Demo.twolevel$cluster)

# Example 6c: Alternative specification without using the '...' argument
item.cfa(Demo.twolevel[, c("y4", "y5", "y6", "cluster")], cluster = "cluster")

#-----
# Print argument

# Example 7a: Request all results
item.cfa(HolzingerSwineford1939, x1, x2, x3, print = "all")

# Example 7b: Request modification indices with value equal or higher than 5
item.cfa(HolzingerSwineford1939, x1, x2, x3, x4, print = "modind", mod.minval = 5)

#-----
# lavaan summary of the estimated model

# Example 8
mod <- item.cfa(HolzingerSwineford1939, x1, x2, x3, output = FALSE)

lavaan::summary(mod$model.fit, standardized = TRUE, fit.measures = TRUE)

```

```
## Not run:
#-----
# Write Results

# Example 9a: Write Results into a text file
item.cfa(HolzingerSwineford1939, x1, x2, x3, write = "CFA.txt")

# Example 9b: Write Results into a Excel file
item.cfa(HolzingerSwineford1939, x1, x2, x3, write = "CFA.xlsx")

## End(Not run)
```

item.invar	<i>Between-Group and Longitudinal Measurement Invariance Evaluation</i>
------------	---

Description

This function evaluates configural, metric, scalar, and strict between-group or longitudinal (partial) measurement invariance using confirmatory factor analysis with continuous indicators by calling the `cfa` function in the R package **lavaan**. By default, the function evaluates configural, metric, and scalar measurement invariance by providing a table with model fit information (i.e., chi-square test, fit indices based on a proper null model, and information criteria) and model comparison (i.e., chi-square difference test, change in fit indices, and change in information criteria). Additionally, variance-covariance coverage of the data, descriptive statistics, parameter estimates, modification indices, and residual correlation matrix can be requested by specifying the argument `print`.

Usage

```
item.invar(data, ..., model = NULL, rescov = NULL, rescov.long = TRUE,
  group = NULL, long = FALSE, cluster = NULL,
  invar = c("config", "metric", "scalar", "strict"),
  partial = NULL, ident = c("marker", "var", "effect"),
  estimator = c("ML", "MLM", "MLMV", "MLMVS", "MLF", "MLR",
    "GLS", "WLS", "DWLS", "WLSM", "WLSMV",
    "ULS", "ULSM", "ULSMV", "DLS", "PML"),
  missing = c("listwise", "pairwise", "fiml", "two.stage",
    "robust.two.stage", "doubly.robust"), null.model = TRUE,
  print = c("all", "summary", "coverage", "descript", "fit", "est",
    "modind", "resid"),
  print.fit = c("all", "standard", "scaled", "robust"),
  mod.minval = 6.63, resid.minval = 0.1, digits = 3, p.digits = 3,
  as.na = NULL, write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a data frame. If <code>model = NULL</code> , confirmatory factor analysis based on a measurement model with one factor labeled <code>f</code> comprising all variables in the data frame specified in <code>x</code> for evaluating between-group measurement invariance for the grouping variable specified in the argument <code>group</code> is conducted. Longitudinal measurement invariance evaluation can only be conducted by specifying the model using the argument <code>model</code> . Note that the cluster variable is excluded from <code>x</code> when specifying <code>cluster</code> . If <code>model</code> is specified, the data frame needs to contain all variables used in the argument <code>model</code> and the cluster variable when specifying the name of the cluster variable in the argument <code>cluster</code> .
<code>...</code>	an expression indicating the variable names in <code>data</code> , e.g., <code>item.invar(dat, x1, x2, x2, group = "group")</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>model</code>	a character vector specifying a measurement model with one factor, or a list of character vectors for specifying a measurement model with more than one factor for evaluating between-group measurement invariance when <code>long = FALSE</code> or a list of character vectors for specifying a measurement model with one factor for each time of measurement for evaluating longitudinal measurement invariance when specifying <code>long = TRUE</code> . For example, <code>model = c("x1", "x2", "x3", "x4")</code> for specifying a measurement model with one factor labeled <code>f</code> comprising four indicators, or <code>model = list(factor1 = c("x1", "x2", "x3", "x4"), factor2 = c("x5", "x6", "x7", "x8"))</code> for specifying a measurement model with two latent factors labeled <code>factor1</code> and <code>factor2</code> each comprising four indicators for evaluating between-group measurement invariance, or <code>model = list(time1 = c("ax1", "ax2", "ax3", "ax4"), time2 = c("bx1", "bx2", "bx3", "bx4"), time3 = c("cx1", "cx2", "cx3", "cx4"))</code> for specifying a longitudinal measurement model with three time points comprising four indicators at each time point. This function cannot evaluate longitudinal measurement invariance for a measurement model with more than one factor. Note that the name of each list element is used to label factors, i.e., all list elements need to be named, otherwise factors are labeled with <code>"f1"</code> , <code>"f2"</code> , <code>"f3"</code> when <code>long = FALSE</code> and with <code>"t1"</code> , <code>"t2"</code> , <code>"t3"</code> when <code>long = TRUE</code> and so on.
<code>rescov</code>	a character vector or a list of character vectors for specifying residual covariances, e.g., <code>rescov = c("x1", "x2")</code> for specifying a residual covariance between items <code>x1</code> and <code>x2</code> , or <code>rescov = list(c("x1", "x2"), c("x3", "x4"))</code> for specifying residual covariances between items <code>x1</code> and <code>x2</code> , and items <code>x3</code> and <code>x4</code> .
<code>rescov.long</code>	logical: if <code>TRUE</code> (default), residual covariances between parallel indicators are estimated across time when evaluating longitudinal measurement invariance (<code>long = TRUE</code>), i.e., residual variances of the same indicators that are measured at different time points are correlated across all possible time points. Note that residual covariances should be estimated even if the parameter estimates are statistically not significant since indicator-specific systematic variance is likely to correlate with itself over time (Little, 2013, p. 164).
<code>group</code>	either a character string indicating the variable name of the grouping variable in the data frame specified in <code>x</code> or a vector representing the groups for conducting multiple-group analysis to evaluate between-group measurement invariance.

long	logical: if TRUE, longitudinal measurement invariance evaluation is conducted. The longitudinal measurement model is specified by using the argument <code>model</code> . Note that this function can only evaluate either between-group or longitudinal measurement invariance, but not both at the same time.
cluster	either a character string indicating the variable name of the cluster variable in data, or a vector representing the nested grouping structure (i.e., group or cluster variable) for computing cluster-robust standard errors. Note that cluster-robust standard errors are not available when treating indicators of the measurement model as ordered (ordinal).
invar	a character string indicating the level of measurement invariance to be evaluated, i.e., <code>config</code> to evaluate configural measurement invariance (i.e., same factor structure across groups or time), <code>metric</code> to evaluate configural and metric measurement invariance (i.e., equal factor loadings across groups or time), <code>scalar</code> (default) to evaluate configural, metric and scalar measurement invariance (i.e., equal intercepts or thresholds across groups or time), and <code>strict</code> to evaluate configural, metric, scalar, and strict measurement invariance (i.e., equal residual variances across groups or time).
partial	a character string or character vector containing the labels of the parameters which should be free in all groups or across time to specify a partial measurement invariance model. Note that the labels of the parameters need to match the labels shown in the output, i.e., "L" with a number for factor loadings, "T" with a number for intercepts, and "E" with a number for factor residual variances. The number attached to the "L", "T", or "E" label corresponds to the number of the indicator in the measurement model (e.g., "T3" for the intercept of the third indicator). When specifying the model using the argument <code>model</code> , however, the number for the factor loading is a combination of the number of the factor and the number of the indicator (e.g., "L23" is the third indicator of the second factor). Note that at least two invariant indicators are needed for a partial measurement invariance model. Otherwise there might be issues with model non-identification.
ident	a character string indicating the method used for identifying and scaling latent variables, i.e., <code>"marker"</code> for the marker variable method fixing the first factor loading of each latent variable to 1, <code>"var"</code> (default) for the fixed variance method fixing the variance of each latent variable to 1, or <code>"effect"</code> for the effects-coding method using equality constraints so that the average of the factor loading for each latent variable equals 1.
estimator	a character string indicating the estimator to be used (see 'Details' in the help page of the <code>item.cfa()</code> function). By default, <code>"MLR"</code> is used for CFA models with continuous indicators.
missing	a character string indicating how to deal with missing data, i.e., <code>"listwise"</code> for listwise deletion, <code>"pairwise"</code> for pairwise deletion, <code>"fiml"</code> for full information maximum likelihood method, <code>two.stage</code> for two-stage maximum likelihood method, <code>robust.two.stage</code> for robust two-stage maximum likelihood method, and <code>doubly-robust</code> for doubly-robust method (see 'Details' in the help page of the <code>item.cfa()</code> function). By default, <code>"fiml"</code> is used for CFA models with continuous indicators which are estimated by using <code>estimator = "MLR"</code> . However, argument <code>missing</code> switches to <code>listwise</code> when the data set is complete, i.e., it

is not possible to use FIML in complete data. Note that the robust CFI, TLI, and RMSEA are different in complete data depending on whether FIML or listwise deletion was specified when estimating the model in lavaan.

<code>null.model</code>	logical: if TRUE (default), the proper null model for computing incremental fit indices (i.e., CFI and TLI) is used, i.e., means and variances of the indicators are constrained to be equal across group or time in the null model (Little, 2013, p. 112).
<code>print</code>	a character string or character vector indicating which results to show on the console, i.e. "all" for all results, "summary" for a summary of the specification of the estimation method and missing data handling in lavaan, "coverage" for the variance-covariance coverage of the data, "descript" for descriptive statistics, "fit" for model fit and model comparison, "est" for parameter estimates, "modind" for modification indices, and "resid" for the residual correlation matrix and standardized residual means. By default, a summary of the specification, model fit, and parameter estimates are printed. Note that parameter estimates, modification indices, and residual correlation matrix is only provided for the model investigating the level of measurement invariance specified in the argument "invar".
<code>print.fit</code>	a character string or character vector indicating which version of the CFI, TLI, and RMSEA to show on the console when using a robust estimation method involving a scaling correction factor, i.e., "all" for all versions of the CFI, TLI, and RMSEA, "standard" (default when estimator is one of "ML", "MLF", "GLS", "WLS", "DWLS", "ULS", "PML") for fit indices without any non-normality correction, "scaled" for population-corrected robust fit indices with ad hoc non-normality correction, and robust (default when estimator is one of "MLM", "MLMV", "MLMVS", "MLR", "WLSM", "WLSMV", "ULSM", "ULSMV", "DLS") for sample-corrected robust fit indices based on formula provided by Li and Bentler (2006) and Brosseau-Liard and Savalei (2014).
<code>mod.minval</code>	numeric value to filter modification indices and only show modifications with a modification index value equal or higher than this minimum value. By default, modification indices equal or higher 6.63 are printed. Note that a modification index value of 6.63 is equivalent to a significance level of $\alpha = .01$.
<code>resid.minval</code>	numeric value indicating the minimum absolute residual correlation coefficients and standardized means to highlight in boldface. By default, absolute residual correlation coefficients and standardized means equal or higher 0.1 are highlighted. Note that highlighting can be disabled by setting the minimum value to 1.
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying results. Note that information criteria and chi-square test statistic are printed with <code>digits</code> minus 1 decimal places.
<code>p.digits</code>	an integer value indicating the number of decimal places to be used for displaying <i>p</i> -values, covariance coverage (i.e., <code>p.digits - 1</code>), and residual correlation coefficients.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e., these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to <code>x</code> but not to <code>group</code> or <code>cluster</code> .

write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked and convergence and model identification checks are conducted for all estimated models.
output	logical: if TRUE (default), output is shown.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	data frame including all variables used in the analysis, i.e., indicators for the factor, grouping variable and cluster variable
args	specification of function arguments
model	list with specified model for the configural, metric, scalar, and strict invariance model
model.fit	list with fitted lavaan object of the configural, metric, scalar, and strict invariance model
check	list with the results of the convergence and model identification check for the configural, metric, scalar, and strict invariance model
result	list with result tables, i.e., summary for the summary of the specification of the estimation method and missing data handling in lavaan, coverage for the variance-covariance coverage of the data, <code>descript</code> for descriptive statistics, <code>fit</code> for a list with model fit based on standard, scaled, and robust fit indices, <code>est</code> for a list with parameter estimates for the configural, metric, scalar, and strict invariance model, <code>modind</code> for the list with modification indices for the configural, metric, scalar, and strict invariance model, <code>score</code> for the list with result of the score tests for constrained parameters for the configural, metric, scalar, and strict invariance model, and <code>resid</code> for the list with residual correlation matrices and standardized residual means for the configural, metric, scalar, and strict invariance model

Note

The function uses the functions `cfa`, `fitmeasures`, `lavInspect`, `lavTech`, `lavTestLRT`, `lavTestScore`, `modindices`, `parameterEstimates`, `parTable`, and `standardizedsolution` provided in the R package **lavaan** by Yves Rosseel (2012).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Brosseau-Liard, P. E., & Savalei, V. (2014) Adjusting incremental fit indices for nonnormality. *Multivariate Behavioral Research*, 49, 460-470. <https://doi.org/10.1080/00273171.2014.933697>
- Li, L., & Bentler, P. M. (2006). Robust statistical tests for evaluating the hypothesis of close fit of misspecified mean and covariance structural models. *UCLA Statistics Preprint #506*. University of California.
- Little, T. D. (2013). *Longitudinal structural equation modeling*. Guilford Press.
- Rosseel, Y. (2012). lavaan: An R Package for Structural Equation Modeling. *Journal of Statistical Software*, 48, 1-36. <https://doi.org/10.18637/jss.v048.i02>

See Also

[item.cfa](#), [multilevel.invar](#), [write.result](#)

Examples

```
## Not run:

# Load data set "HolzingerSwineford1939" in the lavaan package
data("HolzingerSwineford1939", package = "lavaan")

#-----
# Between-Group Measurement Invariance Evaluation

#.....
# Measurement model with one factor

# Example 1a: Specification using the argument '...'
item.invar(HolzingerSwineford1939, x1:x4, group = "sex")

# Example 1b: Alternative specification without using the argument '...'
item.invar(HolzingerSwineford1939[, c("x1", "x2", "x3", "x4")],
           group = HolzingerSwineford1939$sex)

# Example 1c: Alternative specification without using the argument '...'
item.invar(HolzingerSwineford1939[, c("x1", "x2", "x3", "x4", "sex")], group = "sex")

# Example 1d: Alternative specification using the argument 'model'
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"), group = "sex")

#.....
# Measurement model with two factors

item.invar(HolzingerSwineford1939,
           model = list(c("x1", "x2", "x3", "x4"), c("x5", "x6", "x7", "x8")),
           group = "sex")

#.....
# Configural, metric, scalar, and strict measurement invariance
```

```

# Example 2: Evaluate configural, metric, scalar, and strict measurement invariance
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", invar = "strict")

#.....
# Partial measurement invariance

# Example 3: Free second factor loading (L2) and third intercept (T3)
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", partial = c("L2", "T3"), print = c("fit", "est"))

#.....
# Residual covariances

# Example 4a: One residual covariance
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           rescov = c("x3", "x4"), group = "sex")

# Example 4b: Two residual covariances
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           rescov = list(c("x1", "x4"), c("x3", "x4")), group = "sex")

#.....
# Scaled test statistic and cluster-robust standard errors

# Example 5a: Specify cluster variable using a variable name in 'data'
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", cluster = "agemo")

# Example 5b: Specify vector of the cluster variable in the argument 'cluster'
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", cluster = HolzingerSwineford1939$agemo)

#.....
# Default Null model

# Example 6: Specify default null model for computing incremental fit indices
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", null.model = FALSE)

#.....
# Print argument

# Example 7a: Request all results
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", print = "all")

# Example 7b: Request fit indices with ad hoc non-normality correction
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", print.fit = "scaled")

# Example 7c: Request modification indices with value equal or higher than 10
# and highlight residual correlations equal or higher than 0.3

```

```

item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", print = c("modind", "resid"),
           mod.minval = 10, resid.minval = 0.3)

#.....
# Model syntax and lavaan summary of the estimated model

# Example 8
mod <- item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
                  group = "sex", output = FALSE)

# lavaan model syntax scalar invariance model
cat(mod$model$scalar)

# lavaan summary of the scalar invariance model
lavaan::summary(mod$model.fit$scalar, standardized = TRUE, fit.measures = TRUE)

#-----
# Longitudinal Measurement Invariance Evaluation

# Example 9: Two time points with three indicators at each time point
item.invar(HolzingerSwineford1939,
           model = list(c("x1", "x2", "x3"),
                        c("x5", "x6", "x7")), long = TRUE)

#-----
# Write Results

# Example 10a: Write Results into a text file
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", print = "all", write = "Invariance.txt", output = FALSE)

# Example 10b: Write Results into a Excel file
item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
           group = "sex", print = "all", write = "Invariance.xlsx", output = FALSE)

result <- item.invar(HolzingerSwineford1939, model = c("x1", "x2", "x3", "x4"),
                    group = "sex", print = "all", output = FALSE)
write.result(result, "Invariance.xlsx")

## End(Not run)

```

item.omega

Coefficient Omega, Hierarchical Omega, and Categorical Omega

Description

This function computes point estimate and confidence interval for the coefficient omega (McDonald, 1978), hierarchical coefficient omega (Kelley & Pornprasertmanit, 2016), and categorical coefficient omega (Green & Yang, 2009) along with standardized factor loadings and omega if item

deleted. By default, the function computes coefficient omega based on maximum likelihood parameter (ML) estimates using full information maximum likelihood (FIML) method in the presence of missing data.

Usage

```
item.omega(data, ..., rescov = NULL, type = c("omega", "hierarch", "categ"),
            exclude = NULL, std = FALSE,
            estimator = c("ML", "GLS", "WLS", "DWLS", "ULS", "PML"),
            missing = c("listwise", "pairwise", "fiml"),
            print = c("all", "omega", "item"), digits = 2, conf.level = 0.95,
            as.na = NULL, write = NULL, append = TRUE, check = TRUE,
            output = TRUE)
```

Arguments

data	a data frame. Note that at least three items are needed for computing coefficient omega
...	an expression indicating the variable names in data e.g., <code>item.omega(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
rescov	a character vector or a list of character vectors for specifying residual covariances when computing coefficient omega, e.g. <code>rescov = c("x1", "x2")</code> for specifying a residual covariance between items x1 and x2 or <code>rescov = list(c("x1", "x2"), c("x3", "x4"))</code> for specifying residual covariances between items x1 and x2, and items x3 and x4.
type	a character string indicating the type of omega to be computed, i.e., <code>omega</code> (default) for coefficient omega, <code>hierarch</code> for hierarchical coefficient omega, and <code>categ</code> for categorical coefficient omega.
exclude	a character vector indicating items to be excluded from the analysis.
std	logical: if TRUE, the standardized coefficient omega is computed.
estimator	a character string indicating the estimator to be used (see 'Details' in the item.cfa function). By default, "ML" is used for computing (hierarchical) coefficient omega and "DWLS" is used for computing ordinal coefficient omega.
missing	a character string indicating how to deal with missing data. (see 'Details' in the item.cfa function). By default, full information maximum likelihood method (<code>missing = "fiml"</code>) is used for computing (hierarchical) coefficient omega and pairwise deletion (<code>missing = "pairwise"</code>) is used to compute coefficient omega.
print	a character vector indicating which results to show, i.e. "all" for all results "omega" (default) for the coefficient omega, and "item" for item statistics.
digits	an integer value indicating the number of decimal places to be used for displaying omega and standardized factor loadings.
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.

write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown.

Details

Coefficient omega is computed by conducting a confirmatory factor analysis based on the congeneric measurement model (Graham, 2006) using the `cfa()` function in the **lavaan** package by Yves Rosseel (2019).

Approximate confidence intervals are computed using the procedure by Feldt, Woodruff and Salih (1987). Note that there are at least 10 other procedures for computing the confidence interval (see Kelley and Pornprasertmanit, 2016), which are implemented in the `ci.reliability()` function in the **MBESS** package by Ken Kelley (2019).

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	data frame used for the current analysis
args	specification of function arguments
model.fit	fitted lavaan object (<code>mod.fit</code>)
result	list with result tables, i.e., omega for a table with coefficient omega and <code>itemstat</code> for a table with item statistics

Note

Computation of the hierarchical and categorical omega is based on the `ci.reliability()` function in the **MBESS** package by Ken Kelley (2019).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Chalmers, R. P. (2018). On misconceptions and the limited usefulness of ordinal alpha. *Educational and Psychological Measurement*, 78, 1056-1071. <https://doi.org/10.1177/0013164417727036>
- Cronbach, L.J. (1951). Coefficient alpha and the internal structure of tests. *Psychometrika*, 16, 297-334. <https://doi.org/10.1007/BF02310555>

- Cronbach, L.J. (2004). My current thoughts on coefficient alpha and successor procedures. *Educational and Psychological Measurement*, 64, 391-418. <https://doi.org/10.1177/0013164404266386>
- Feldt, L. S., Woodruff, D. J., & Salih, F. A. (1987). Statistical inference for coefficient alpha. *Applied Psychological Measurement*, 11 93-103. <https://doi.org/10.1177/014662168701100107>
- Graham, J. M. (2006). Congeneric and (essentially) tau-equivalent estimates of score reliability: What they are and how to use them. *Educational and Psychological Measurement*, 66(6), 930-944. <https://doi.org/10.1177/0013164406288165>
- Kelley, K., & Pornprasertmanit, S. (2016). Confidence intervals for population reliability coefficients: Evaluation of methods, recommendations, and software for composite measures. *Psychological Methods*, 21, 69-92. <https://doi.org/10.1037/a0040086>.
- Ken Kelley (2019). *MBESS: The MBESS R Package*. R package version 4.6.0. <https://CRAN.R-project.org/package=MBESS>
- Revelle, W. (2025). *psych: Procedures for psychological, psychometric, and personality research*. Northwestern University, Evanston, Illinois. R package version 2.5.3, <https://CRAN.R-project.org/package=psych>.
- Zumbo, B. D., & Kroc, E. (2019). A measurement is a choice and Stevens' scales of measurement do not help make it: A response to Chalmers. *Educational and Psychological Measurement*, 79, 1184-1197. <https://doi.org/10.1177/0013164419844305>
- Zumbo, B. D., Gadermann, A. M., & Zeisser, C. (2007). Ordinal versions of coefficients alpha and theta for Likert rating scales. *Journal of Modern Applied Statistical Methods*, 6, 21-29. <https://doi.org/10.22237/jmasm/1177992180>

See Also

[item.omega](#), [item.cfa](#), [item.invar](#), [item.reverse](#), [item.scores](#), [write.result](#)

Examples

```
## Not run:
dat <- data.frame(item1 = c(3, NA, 3, 4, 1, 2, 4, 2), item2 = c(5, 3, 3, 2, 2, 1, 3, 1),
                  item3 = c(4, 2, 4, 2, 1, 3, 4, 1), item4 = c(4, 1, 2, 2, 1, 3, 4, 3))

# Example 1a: Coefficient omega, full information maximum likelihood method
item.omega(dat)

# Example 1b: Coefficient omega, listwise deletion
item.omega(dat, missing = "listwise")

# Example 2: Coefficient omega and item statistics after excluding item3
item.omega(dat, exclude = "item3", print = "all")

# Example 3a: Coefficient omega with a residual covariance
item.omega(dat, rescov = c("item1", "item2"))

# Example 3b: Coefficient omega with residual covariances
item.omega(dat, rescov = list(c("item1", "item2"), c("item1", "item3"))))

# Example 4: Ordinal coefficient omega and item statistics
item.omega(dat, type = "categ", print = "all")
```

```
# Example 6: Summary of the CFA model used to compute coefficient omega
lavaan::summary(item.omega(dat, output = FALSE)$model.fit,
  fit.measures = TRUE, standardized = TRUE)

# Example 7a: Write Results into a text file
item.omega(dat, write = "Omega.txt")

# Example 7b: Write Results into a Excel file
item.omega(dat, write = "Omega.xlsx")

## End(Not run)
```

item.reverse

Reverse Code Scale Item

Description

This function reverse codes inverted items, i.e., items that are negatively worded.

Usage

```
item.reverse(data, ..., min = NULL, max = NULL, keep = NULL, append = TRUE,
  name = ".r", as.na = NULL, table = FALSE, check = TRUE)
```

Arguments

data	a numeric vector for reverse coding an item or data frame for reverse coding more than one item.
...	an expression indicating the variable names in data e.g., <code>item.reverse(x1, x2, x3, data = dat)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
min	an integer indicating the minimum of the item (i.e., lowest possible scale value).
max	an integer indicating the maximum of the item (i.e., highest possible scale value).
keep	a numeric vector indicating values not to be reverse coded.
append	logical: if TRUE (default), recoded variable(s) are appended to the data frame specified in the argument data.
name	a character string or character vector indicating the names of the reverse coded item. By default, variables are named with the ending <code>".r"</code> resulting in e.g. <code>"x1.r"</code> and <code>"x2.r"</code> . Variable names can also be specified using a character vector matching the number of variables (e.g., <code>name = c("reverse.x1", "reverse.x2")</code>).
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
table	logical: if TRUE, a cross table item x reverse coded item is printed on the console if only one variable is specified.
check	logical: if TRUE (default), argument specification is checked.

item.scores

*Scale Scores***Description**

This function computes (prorated) scale scores by averaging the (available) items that measure a single construct by default.

Usage

```
item.scores(data, ..., fun = c("mean", "sum", "median", "var", "sd", "min", "max"),
            prorated = TRUE, p.avail = NULL, n.avail = NULL, append = TRUE,
            name = "scores", as.na = NULL, check = TRUE)
```

Arguments

data	a data frame with numeric vectors.
...	an expression indicating the variable names in data, e.g., <code>item.scores(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
fun	a character string indicating the function used to compute scale scores, default: "mean".
prorated	logical: if TRUE (default), prorated scale scores are computed (see 'Details'); if FALSE, scale scores of only complete cases are computed.
p.avail	a numeric value indicating the minimum proportion of available item responses needed for computing a prorated scale score for each case, e.g. <code>p.avail = 0.8</code> indicates that scale scores are only computed for cases with at least 80% of item responses available. By default prorated scale scores are computed for all cases with at least one item response. Note that either argument <code>p.avail</code> or <code>n.avail</code> is used to specify the proration criterion.
n.avail	an integer indicating the minimum number of available item responses needed for computing a prorated scale score for each case, e.g. <code>n.avail = 2</code> indicates that scale scores are only computed for cases with item responses on at least 2 items. By default prorated scale scores are computed for all cases with at least one item response. Note that either argument <code>p.avail</code> or <code>n.avail</code> is used to specify the proration criterion.
append	logical: if TRUE (default), a variable with scale scores is appended to the data frame specified in the argument <code>data</code> .
name	a character string indicating the names of the variable appended to the data frame specified in the argument <code>data</code> when <code>append = TRUE</code> . By default, the variable is named <code>scores</code> .
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
check	logical: if TRUE (default), argument specification is checked.

Details

Prorated mean scale scores are computed by averaging the available items, e.g., if a participant answers 4 out of 8 items, the prorated scale score is the average of the 4 responses. Averaging the available items is equivalent to substituting the mean of a participant's own observed items for each of the participant's missing items, i.e., *person mean imputation* (Mazza, Enders & Ruehlman, 2015) or *ipsative mean imputation* (Schafer & Graham, 2002).

Proration may be reasonable when (1) a relatively high proportion of the items (e.g., 0.8) and never fewer than half are used to form the scale score, (2) means of the items comprising a scale are similar and (3) the item-total correlations are similar (Enders, 2010; Graham, 2009; Graham, 2012). Results of simulation studies indicate that proration is prone to substantial bias when either the item means or the inter-item correlation vary (Lee, Bartholow, McCarthy, Pederson & Sher, 2014; Mazza et al., 2015).

Value

Returns a numeric vector with the same length as `nrow(x)` containing (prorated) scale scores.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Enders, C. K. (2010). *Applied missing data analysis*. New York, NY: Guilford Press.
- Graham, J. W. (2009). Missing data analysis: Making it work in the real world. *Annual Review of Psychology*, 60, 549-576. <https://doi.org/10.1146/annurev.psych.58.110405.085530>
- Graham, J. W. (2012). *Missing data: Analysis and design*. New York, NY: Springer
- Lee, M. R., Bartholow, B. D., McCarthy, D. M., Pederson, S. L., & Sher, K. J. (2014). Two alternative approaches to conventional person-mean imputation scoring of the self-rating of the effects of alcohol scale (SRE). *Psychology of Addictive Behaviors*, 29, 231-236. <https://doi.org/10.1037/adb0000015>
- Mazza, G. L., Enders, C. G., & Ruehlman, L. S. (2015). Addressing item-level missing data: A comparison of proration and full information maximum likelihood estimation. *Multivariate Behavioral Research*, 50, 504-519. <https://doi.org/10.1080/00273171.2015.1068157>
- Schafer, J. L., & Graham, J. W. (2002). Missing data: Our view of the state of the art. *Psychological Methods*, 7, 147-177. <https://doi.org/10.1037/1082-989X.7.2.147>

See Also

[cluster.scores](#), [item.alpha](#), [item.cfa](#), [item.omega](#),

Examples

```
dat <- data.frame(item1 = c(3, 2, 4, 1, 5, 1, 3, NA),
  item2 = c(2, 2, NA, 2, 4, 2, NA, 1),
  item3 = c(1, 1, 2, 2, 4, 3, NA, NA),
  item4 = c(4, 2, 4, 4, NA, 2, NA, NA),
  item5 = c(3, NA, NA, 2, 4, 3, NA, 3))
```

```
# Example 1: Prorated mean scale scores
item.scores(dat)

# Example 2: Prorated standard deviation scale scores
item.scores(dat, fun = "sd")

# Example 3: Sum scale scores without proration
item.scores(dat, fun = "sum", prorated = FALSE)

# Example 4: Prorated mean scale scores,
# minimum proportion of available item responses = 0.8
item.scores(dat, p.avail = 0.8)

# Example 5: Prorated mean scale scores,
# minimum number of available item responses = 3
item.scores(dat, n.avail = 3)
```

lagged	Create Lagged Variables
--------	-------------------------

Description

This function computes lagged values of variables by a specified number of observations. By default, the function returns lag-1 values of the vector or data frame specified in the first argument.

Usage

```
lagged(data, ..., id = NULL, obs = NULL, day = NULL, lag = 1, time = NULL,
       units = c("secs", "mins", "hours", "days", "weeks"), append = TRUE,
       name = ".lag", name.td = ".td", as.na = NULL, check = TRUE)
```

Arguments

data	a numeric vector for computing a lagged values for a variable or data frame for computing lagged values for more than one variable. Note that the subject ID variable (id), observation number variable (obs), day number variable (day), and the date and time variable (time) are excluded from data when specifying theses arguments.
...	an expression indicating the variable names in data. Note that the operators +, -, ~, :, ::, and ! can also be used to select variables, see 'Details' in the df.subset function.
id	either a character string indicating the variable name of the subject ID variable or a vector representing the subject IDs, see 'Details'.
obs	either a character string indicating the variable name of the observation number variable or a vector representing the observations. Note that duplicated values within the same subject ID are not allowed, see 'Details'.
day	either a character string indicating the variable name of the day number variable in or a vector representing the days, see 'Details'.

lag	a numeric value specifying the lag, e.g. lag = 1 (default) returns lag-1 values.
time	a variable of class POSIXct or POSIXlt representing the date and time of the observation used to compute time differences between observations.
units	a character string indicating the units in which the time difference is represented, i.e., "secs" for seconds, "mins" (default) for minutes, "hours" for hours, "days" for days, and "weeks" for weeks.
append	logical: if TRUE (default), lagged variable(s) are appended to the data frame specified in the argument data.
name	a character string or character vector indicating the names of the lagged variables. By default, lagged variables are named with the ending ".lag" resulting in e.g. "x1.lag" and "x2.lag" when specifying two variables. Variable names can also be specified using a character vector matching the number of variables, e.g., name = c("lag.x1", "lag.x2").
name.td	a character string or character vector indicating the names of the time difference variables when specifying a date and time variables for the argument time. By default, time difference variables are named with the ending ".td" resulting in e.g. "x1.td" and "x2.td" when specifying two variables. Variable names can also be specified using a character vector matching the number of variables specified, e.g., name = c("td.x1", "td.x2").
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that as.na() function is only applied to the argument data, but not to cluster.
check	logical: if TRUE (default), argument specification is checked.

Details

The function is used to create lagged version of the variable(s) specified via the data argument:

If the `id` argument is not specified i.e., `id = NULL`, all observations are assumed to come from the same subject. If the dataset includes multiple subjects, then this variable needs to be specified so that observations are not lagged across subjects

Optional argument `day` If the `day` argument is not specified i.e., `day = NULL`, values of the variable to be lagged are allowed to be lagged across days in case there are multiple observation days.

Optional argument `obs` If the `obs` argument is not specified i.e., `obs = NULL`, consecutive observations from the same subjects are assumed to be one lag apart.

Value

Returns a numeric vector or data frame with the same length or same number of rows as data containing the lagged variable(s).

Note

This function is based on the `lagvar()` function in the **esmpack** package by Wolfgang Viechtbauer and Mihail Constantin (2023).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Viechtbauer W, Constantin M (2023). *esmpack: Functions that facilitate preparation and management of ESM/EMA data*. R package version 0.1-20.

See Also

[center](#), [rec](#), [coding](#), [item.reverse](#).

Examples

```
dat <- data.frame(subject = rep(1:2, each = 6),
                  day = rep(1:2, each = 3),
                  obs = rep(1:6, times = 2),
                  time = as.POSIXct(c("2024-01-01 09:01:00", "2024-01-01 12:05:00",
                                     "2024-01-01 15:14:00", "2024-01-02 09:03:00",
                                     "2024-01-02 12:21:00", "2024-01-02 15:03:00",
                                     "2024-01-01 09:02:00", "2024-01-01 12:09:00",
                                     "2024-01-01 15:06:00", "2024-01-02 09:02:00",
                                     "2024-01-02 12:15:00", "2024-01-02 15:06:00")),
                  pos = c(6, 7, 5, 8, NA, 7, 4, NA, 5, 4, 5, 3),
                  neg = c(2, 3, 2, 5, 3, 4, 6, 4, 6, 4, NA, 8))

# Example 1: Lagged variable for 'pos'
lagged(dat$pos, id = dat$subject, day = dat$day)

# Example 1b: Alternative specification without using the '...' argument
lagged(dat[, c("pos", "subject", "day")], id = "subject", day = "day")

# Example 1c: Alternative specification using the 'data' argument
lagged(pos, data = dat, id = "subject", day = "day")

# Example 2a: Lagged variable for 'pos' and 'neg'
lagged(dat[, c("pos", "neg")], id = dat$subject, day = dat$day)

# Example 2b: Alternative specification using the 'data' argument
lagged(pos, neg, data = dat, id = "subject", day = "day")

# Example 3: Lag-2 variables for 'pos' and 'neg'
lagged(pos, neg, data = dat, id = "subject", day = "day", lag = 2)

# Example 4: Lagged variable and time difference variable
lagged(pos, neg, data = dat, id = "subject", day = "day", time = "time")

# Example 5: Lagged variables and time difference variables,
# name variables
lagged(pos, neg, data = dat, id = "subject", day = "day", time = "time",
       name = c("p.lag1", "n.lag1"), name.td = c("p.diff", "n.diff"))
```

```
# Example 6: NA observations excluded from the data frame
dat.excl <- dat[!is.na(dat$pos), ]

# Number of observation not taken into account, i.e.,
# - observation 4 used as lagged value for observation 6 for subject 1
# - observation 1 used as lagged value for observation 3 for subject 2
lagged(pos, data = dat.excl, id = "subject", day = "day")

# Number of observation taken into account by specifying the 'ob' argument
lagged(pos, data = dat.excl, id = "subject", day = "day", obs = "obs")
```

libraries
Load and Attach Multiple Packages

Description

This function loads and attaches multiple add-on packages at once.

Usage

```
libraries(..., install = FALSE, quiet = TRUE, check = TRUE, output = TRUE)
```

Arguments

...	the names of the packages to be loaded, given as names (e.g., <code>misty</code> , <code>lavaan</code> , <code>lme4</code>), or literal character strings (e.g., <code>"misty"</code> , <code>"lavaan"</code> , <code>"lme4"</code>), or character vector (e.g., <code>c("misty", "lavaan", "lme4")</code>).
<code>install</code>	logical: if TRUE, missing packages and dependencies are installed.
<code>quiet</code>	logical: if TRUE (default), startup messages when loading package are disabled.
<code>check</code>	logical: if TRUE, argument specification is checked.
<code>output</code>	logical: logical: if TRUE, output is shown on the console.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Becker, R. A., Chambers, J. M. and Wilks, A. R. (1988) *The New S Language*. Wadsworth & Brooks/Cole.

See Also

[library](#), [require](#)

Examples

```
## Not run:

# Example 1: Load packages using the names of the packages
misty::libraries(misty, lme4, lmerTest)

# Example 2: Load packages using literal character strings
misty::libraries("misty", "lme4", "lmerTest")

# Example 3: Load packages using a character vector
misty::libraries(c("misty", "lme4", "lmerTest"))

# Example 4: Check packages, i.e., TRUE = all depends/imports/suggests installed
misty::libraries(misty, lme4, lmerTest, output = FALSE)$result$restab

# Example 5: Depends, FALSE = not installed, TRUE = installed
misty::libraries(misty, lme4, lmerTest, output = FALSE)$result$depends

# Example 6: Imports, FALSE = not installed, TRUE = installed
misty::libraries(misty, lme4, lmerTest, output = FALSE)$result$imports

# Example 6: Suggests, FALSE = not installed, TRUE = installed
misty::libraries(misty, lme4, lmerTest, output = FALSE)$result$suggests

## End(Not run)
```

mplus

Create, Run, and Print Mplus Models

Description

This wrapper function creates a Mplus input file, runs the input file by using the `mplus.run()` function, and prints the Mplus output file by using the `mplus.print()` function.

Usage

```
mplus(x, file = "Mplus_Input.inp", data = NULL, comment = FALSE,
      replace.inp = TRUE, mplus.run = TRUE, show.out = FALSE,
      replace.out = c("always", "never", "modified"), Mplus = .detect.mplus(),
      print = c("all", "input", "result"),
      input = c("all", "default", "data", "variable", "define", "analysis",
                "model", "montecarlo", "mod.pop", "mod.cov", "mod.miss",
                "message"),
      result = c("all", "default", "summary.analysis.short",
                 "summary.data.short", "random.starts", "summary.fit",
                 "mod.est", "fit", "class.count", "classif", "mod.result",
                 "total.indirect"),
      exclude = NULL, variable = FALSE, not.input = TRUE, not.result = TRUE,
      write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>x</code>	a character string containing the Mplus input text.
<code>file</code>	a character string indicating the name of the Mplus input file with or without the file extension <code>.inp</code> , e.g., <code>"Mplus_Input.inp"</code> or <code>"Mplus_Input"</code> .
<code>data</code>	a matrix or data frame from which the variables names for the subsection NAMES are extracted.
<code>comment</code>	logical: if FALSE (default), comments (i.e., text after the <code>!</code> symbol) are removed from the input text specified in the argument <code>x</code> .
<code>replace.inp</code>	logical: if TRUE (default), an existing input file will be replaced.
<code>mplus.run</code>	logical: if TRUE, the input file specified in the argument <code>file</code> containing the input text specified in the argument <code>x</code> is run using the <code>mplus.run()</code> function.
<code>show.out</code>	logical: if TRUE, estimation output (TECH8) is show on the R console. Note that if run within Rgui, output will display within R, but if run via Rterm, a separate window will appear during estimation.
<code>replace.out</code>	a character string for specifying three settings: <code>"always"</code> (default), which runs all models, regardless of whether an output file for the model exists, <code>"never"</code> , which does not run any model that has an existing output file, and <code>"modified"</code> , which only runs a model if the modified date for the input file is more recent than the output file modified date.
<code>Mplus</code>	a character string for specifying the name or path of the Mplus executable to be used for running models. This covers situations where Mplus is not in the system's path, or where one wants to test different versions of the Mplus program. Note that there is no need to specify this argument for most users since it has intelligent defaults.
<code>print</code>	a character vector indicating which results to show, i.e. <code>"all"</code> (default) for all results <code>"input"</code> for input command sections, and <code>"result"</code> for result sections.
<code>input</code>	a character vector specifying Mplus input command sections included in the output (see 'Details' in the mplus.print function).
<code>result</code>	a character vector specifying Mplus result sections included in the output (see 'Details' in the mplus.print function).
<code>exclude</code>	a character vector specifying Mplus input command or result sections excluded from the output (see 'Details' in the mplus.print function).
<code>variable</code>	logical: if TRUE, names of the variables in the data set (NAMES ARE) specified in the VARIABLE: command section are shown. By default, names of the variables in the data set are excluded from the output unless all variables are used in the analysis (i.e., no USEVARIABLES option specified in the Mplus input file).
<code>not.input</code>	logical: if TRUE (default), character vector indicating the input commands not requested are shown on the console.
<code>not.result</code>	logical: if TRUE (default), character vector indicating the result sections not requested are shown on the console.
<code>write</code>	a character string naming a file for writing the output into a text file with file extension <code>".txt"</code> (e.g., <code>"Output.txt"</code>).

append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console by using the function mplus.print().

Details

The NAMES Option The NAMES option in the VARIABLE section used to assign names to the variables in the data set can be specified by using the data argument:

- Write Mplus Data File: In the first step, the Mplus data file is written by using the write.mplus() function, e.g. write.mplus(ex3_1, file = "ex3_1.dat").
- Specify Mplus Input: In the second step, the Mplus input is specified as a character string. The NAMES option is left out from the Mplus input text, e.g., input <- 'DATA: FILE IS ex3_1.dat;\nMODEL: y1 ON x1 x3; '.
- Run Mplus Input: In the third step, the Mplus input is run by using the mplus() function. The argument data needs to be specified given that the NAMES option was left out from the Mplus input text in the previous step, e.g., mplus(input, file = "ex3_1.inp", data = ex3_1).

Value

Returns an object of class misty.object, which is a list with following entries:

call	function call
type	type of analysis
x	a character vector containing the Mplus input text
args	specification of function arguments
input	list with input command sections
write	write command sections
result	list with input command sections (input) and result sections (result)

Author(s)

Takuya Yanagida

References

Muthen, L. K., & Muthen, B. O. (1998-2017). *Mplus User's Guide* (8th ed.). Muthen & Muthen.

See Also

[read.mplus](#), [write.mplus](#), [mplus.update](#), [mplus.print](#), [mplus.plot](#), [mplus.bayes](#), [mplus.run](#), [mplus.lca](#)

Examples

```
## Not run:

#-----
# Example 1: Write data, specify input, and run input

# Write Mplus Data File
write.mplus(ex3_1, file = "ex3_1.dat")

# Specify Mplus input, specify NAMES option
input1 <- '
DATA:      FILE IS ex3_1.dat;
VARIABLE:  NAMES ARE y1 x1 x3;
MODEL:     y1 ON x1 x3;
OUTPUT:    SAMPSTAT;
'

# Run Mplus input
mplus(input1, file = "ex3_1.inp")

#-----
# Example 2: Alternative specification using the data argument

# Specify Mplus input, leave out the NAMES option
input2 <- '
DATA:      FILE IS ex3_1.dat;
MODEL:     y1 ON x1 x3;
OUTPUT:    SAMPSTAT;
'

# Run Mplus input, specify the data argument
mplus(input2, file = "ex3_1.inp", data = ex3_1)

## End(Not run)
```

mplus.bayes

Mplus Summary Measures, Convergence and Efficiency Diagnostics

Description

This function uses the `h5file` function in the **hdf5r** package to read a Mplus GH5 file that is requested by the command `PLOT: TYPE IS PLOT2` in Mplus to compute point estimates (i.e., mean, median, and MAP), measures of dispersion (i.e., standard deviation and mean absolute deviation), measures of shape (i.e., skewness and kurtosis), credible intervals (i.e., equal-tailed intervals and highest density interval), convergence and efficiency diagnostics (i.e., potential scale reduction factor R -hat, effective sample size, and Monte Carlo standard error), probability of direction, and probability of being in the region of practical equivalence for the posterior distribution for each parameter. By default, the function computes the maximum of rank-normalized split- R -hat and rank normalized folded-split- R -hat, Bulk effective sample size (Bulk-ESS) for rank-normalized values

using split chains, tail effective sample size (Tail-ESS) defined as the minimum of the effective sample size for 0.025 and 0.975 quantiles, the Bulk Monte Carlo standard error (Bulk-MCSE) for the median and Tail Monte Carlo standard error (Tail-MCSE) defined as the maximum of the MCSE for 0.025 and 0.975 quantiles.

Usage

```
mplus.bayes(x,
  print = c("all", "default", "m", "med", "map", "sd", "mad",
    "skew", "kurt", "eti", "hdi",
    "rhat", "b.ess", "t.ess", "b.mcse", "t.mcse"),
  param = c("all", "on", "by", "with", "inter", "var", "r2", "new"),
  std = c("all", "none", "stdyx", "stdy", "std"),
  m.bulk = FALSE, split = TRUE, rank = TRUE, fold = TRUE,
  pd = FALSE, null = 0, rope = NULL,
  ess.tail = c(0.025, 0.975), mcse.tail = c(0.025, 0.975),
  alternative = c("two.sided", "less", "greater"),
  conf.level = 0.95, digits = 2, r.digits = 3, ess.digits = 0,
  mcse.digits = 3, p.digits = 3, write = NULL, append = TRUE,
  check = TRUE, output = TRUE)
```

Arguments

- | | |
|-------|---|
| x | a character string indicating the name of the Mplus GH5 file (HDF5 format) with or without the file extension .gh5, e.g., "Mplus_Plot.gh5" or "Mplus_Plot". |
| print | a character vector indicating which summary measures, convergence, and efficiency diagnostics to be printed on the console, i.e. "all" for all summary measures, convergence, and efficiency diagnostics, "m" for the mean, "med" for the median, "MAP" for the maximum a posteriori probability estimate, "med" for the standard deviation, "mad" for the mean absolute deviation, "skew" for the skewness, "kurt" for the kurtosis, "eti" for the equal-tailed credible interval, "hdi" for the highest density credible interval, "rhat" for the potential scale reduction (PSR) factor R-hat convergence diagnostic, "b.ess" for the bulk effective sample size (ESS), "t.ess" for the tail ESS, "b.mcse" for the bulk Monte Carlo standard error (MCSE), and "t.mcse" for the tail MCSE. The default setting is <code>print = c("med", "sd", "skew", "kurt", "eti", "rhat", "b.ess", "t.ess", "b.mcse", "t.mcse")</code> . |
| param | character vector indicating which parameters to print for the summary measures, convergence, and efficiency diagnostics, i.e., "all" for all parameters, "on" (default), for regression slopes, "by" for factor loadings, "with" for covariances, "inter" for intercepts and thresholds, "var" for (residual) variances, "r2" for r-square, and "new" for parameters not in the analysis model specified in the NEW option. The default setting is "on" if regression slopes are available. Otherwise, the default setting switches to "by" and to "with" if factor loadings are not available. |
| std | a character vector indicating the standardized parameters to print for the summary measures, convergence, and efficiency diagnostics, i.e., "all" for all standardized parameters, "none" (default) for not printing any standardized param- |

	eters, "stdyx" for StdYX standardized parameters, "stdy" for StdY standardized parameters, and "std" for StdX standardized parameters.
<code>m.bulk</code>	logical: if TRUE the Monte Carlo standard error for the mean is computed. The default setting is <code>m.bulk = FALSE</code> , i.e., the Monte Carlo standard error for the median is computed.
<code>split</code>	logical: if TRUE (default), each MCMC chain is split in half before computing R-hat. Note that the argument <code>split</code> is always set to FALSE when computing ESS.
<code>rank</code>	logical: if TRUE (default), rank-normalization is applied to the posterior draws before computing R-hat and ESS. Note that the argument <code>rank</code> is always set to FALSE when computing MCSE.
<code>fold</code>	logical: if TRUE (default), the maximum of rank-normalized split-R-hat and rank normalized folded-split-R-hat is computed. Note that the arguments <code>split</code> and <code>rank</code> are always set to TRUE when specifying <code>fold = TRUE</code> .
<code>pd</code>	logical: if TRUE, the probability of direction is printed on the console.
<code>null</code>	a numeric value considered as a null effect for the probability of direction (default is 0). Note that the value specified in the argument <code>null</code> applies to all parameters which might not be sensible for all parameters.
<code>rope</code>	a numeric vector with two elements indicating the ROPE's lower and upper bounds. ROPE is also depending on the argument <code>alternative</code> , e.g., if <code>rope = c(-0.1, 0.1)</code> , then the actual ROPE is $[-0.1, 0.1]$ given <code>alternative = "two.sided"</code> (default), $[-\text{Inf}, 0.1]$ given <code>alternative = "greater"</code> , and $[-0.1, \text{Inf}]$ given <code>alternative = "less"</code> . Note that the interval specified in the argument <code>rope</code> applies to all parameters which might not be sensible for all parameters.
<code>ess.tail</code>	a numeric vector with two elements to specify the quantiles for computing the tail ESS. The default setting is <code>tail = c(0.025, 0.975)</code> , i.e., tail ESS is the minimum of effective sample sizes for 5% and 95% quantiles.
<code>mcse.tail</code>	a numeric vector with two elements to specify the quantiles for computing the tail MCSE. The default setting is <code>tail = c(0.025, 0.975)</code> , i.e., tail MCSE is the maximum of Monte Carlo standard error for 5% and 95% quantiles.
<code>alternative</code>	a character string specifying the alternative hypothesis for the credible intervals, must be one of "two.sided" (default), "greater" or "less".
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the credible interval. The default setting is <code>conf.level = 0.95</code> .
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying point estimates, measures of dispersion, and credible intervals.
<code>r.digits</code>	an integer value indicating the number of decimal places to be used for displaying R-hat values.
<code>ess.digits</code>	an integer value indicating the number of decimal places to be used for displaying effective sample sizes.
<code>mcse.digits</code>	an integer value indicating the number of decimal places to be used for displaying Monte Carlo standard errors.

p.digits	an integer value indicating the number of decimal places to be used for displaying the probability of direction and the probability of being in the region of practical equivalence (ROPE).
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console by using the function mplus.print().

Details

Convergence and Efficiency Diagnostics for Markov Chains Convergence and efficiency diagnostics for Markov chains is based on following numeric measures:

- **Potential Scale Reduction (PSR) factor R-hat:** The PSR factor R-hat compares the between- and within-chain variance for a model parameter, i.e., R-hat larger than 1 indicates that the between-chain variance is greater than the within-chain variance and chains have not mixed well. According to the default setting, the function computes the improved R-hat as recommended by Vehtari et al. (2020) based on rank-normalizing (i.e., rank = TRUE) and folding (i.e., fold = TRUE) the posterior draws after splitting each MCMC chain in half (i.e., split = TRUE). The traditional R-hat used in Mplus can be requested by specifying split = FALSE, rank = FALSE, and fold = FALSE. Note that the traditional R-hat can catch many problems of poor convergence, but fails if the chains have different variances with the same mean parameter or if the chains have infinite variance with one of the chains having a different location parameter to the others (Vehtari et al., 2020). According to Gelman et al. (2014) a R-hat value of 1.1 or smaller for all parameters can be considered evidence for convergence. The Stan Development Team (2024) recommends running at least four chains and a convergence criterion of less than 1.05 for the maximum of rank normalized split-R-hat and rank normalized folded-split-R-hat. Vehtari et al. (2020), however, recommended to only use the posterior samples if R-hat is less than 1.01 because the R-hat can fall below 1.1 well before convergence in some scenarios (Brooks & Gelman, 1998; Vats & Knudon, 2018).
- **Effective Sample Size (ESS):** The ESS is the estimated number of independent samples from the posterior distribution that would lead to the same precision as the autocorrelated samples at hand. According to the default setting, the function computes the ESS based on rank-normalized split-R-hat and within-chain autocorrelation. The function provides the estimated Bulk-ESS (B.ESS) and the Tail-ESS (T.ESS). The Bulk-ESS is a useful measure for sampling efficiency in the bulk of the distribution (i.e., efficiency of the posterior mean), and the Tail-ESS is useful measure for sampling efficiency in the tails of the distribution (e.g., efficiency of tail quantile estimates). Note that by default, the Tail-ESS is the minimum of the effective sample sizes for 5% and 95% quantiles (tail = c(0.025, 0.975)). According to Kruschke (2015), a rank-normalized ESS greater than 400 is usually sufficient to get a stable estimate of the Monte Carlo standard error. However, a ESS of at least 1000 is considered optimal (Zitzmann & Hecht, 2019).

- **Monte Carlo Standard Error (MCSE):** The MCSE is defined as the standard deviation of the chains divided by their effective sample size and reflects uncertainty due to the stochastic algorithm of the Markov Chain Monte Carlo method. The function provides the estimated Bulk-MCSE (B.MCSE) for the margin of error when using the MCMC samples to estimate the posterior mean and the Tail-ESS (T.MCSE) for the margin of error when using the MCMC samples for interval estimation.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>x</code>	Mplus GH5 file
<code>args</code>	specification of function arguments
<code>data</code>	three-dimensional array parameter x iteration x chain of the posterior
<code>result</code>	result table with summary measures, convergence, and efficiency diagnostics

Note

This function is a modified copy of functions provided in the **rstan** package by Stan Development Team (2024) and **bayestestR** package by Makowski et al. (2019).

Author(s)

Takuya Yanagida

References

- Brooks, S. P. and Gelman, A. (1998). General Methods for Monitoring Convergence of Iterative Simulations. *Journal of Computational and Graphical Statistics*, 7(4): 434–455. MR1665662.
- Gelman, A., & Rubin, D.B. (1992). Inference from iterative simulation using multiple sequences. *Statistical Science*, 7, 457–472. <https://doi.org/10.1214/ss/1177011136>
- Kruschke, J. (2015). *Doing Bayesian data analysis: A tutorial with R, JAGS, and Stan*. Academic Press.
- Makowski, D., Ben-Shachar, M., & Lüdtke, D. (2019). bayestestR: Describing effects and their uncertainty, existence and significance within the Bayesian framework. *Journal of Open Source Software*, 4(40), 1541. <https://doi.org/10.21105/joss.01541>
- Stan Development Team (2024). *RStan: the R interface to Stan*. R package version 2.32.6. <https://mc-stan.org/>.
- Vats, D. and Knudson, C. (2018). Revisiting the Gelman-Rubin Diagnostic. arXiv:1812.09384.
- Vehtari, A., Gelman, A., Simpson, D., Carpenter, B., & Bürkner, P.-C. (2020). Rank-normalization, folding, and localization: An improved R-hat for assessing convergence of MCMC. *Bayesian analysis*, 16(2), 667–718. <https://doi.org/10.1214/20-BA1221>
- Zitzmann, S., & Hecht, M. (2019). Going beyond convergence in Bayesian estimation: Why precision matters too and how to assess it. *Structural Equation Modeling: A Multidisciplinary Journal*, 26(4), 646–661. <https://doi.org/10.1080/10705511.2018.1545232>

See Also

`read.mplus, write.mplus, mplus, mplus.update, mplus.print, mplus.plot, mplus.run, mplus.lca`

Examples

```
## Not run:

#-----
# Mplus Example 3.18: Moderated Mediation with a Plot of the Indirect Effect

# Example 1: Default setting
mplus.bayes("ex3.18.gh5")

# Example 2: Print all parameters
mplus.bayes("ex3.18.gh5", param = "all")

# Example 3: Print parameters not in the analysis model
mplus.bayes("ex3.18.gh5", param = "new")

# Example 4a: Print all summary measures, convergence, and efficiency diagnostics
mplus.bayes("ex3.18.gh5", print = "all")

# Example 4a: Print default measures plus MAP
mplus.bayes("ex3.18.gh5", print = c("default", "map"))

# Example 5: Print traditional R-hat in line with Mplus
mplus.bayes("ex3.18.gh5", split = FALSE, rank = FALSE, fold = FALSE)

# Example 6: Print probability of direction and the probability of
# being ROPE [-0.1, 0.1]
mplus.bayes("ex3.18.gh5", pd = TRUE, rope = c(-0.1, 0.1))

# Example 7: Write Results into a text file
mplus.bayes("ex3.18.gh5", write = "Bayes_Summary.txt")

# Example 8b: Write Results into a Excel file
mplus.bayes("ex3.18.gh5", write = "Bayes_Summary.xlsx")

## End(Not run)
```

mplus.lca

Mplus Model Specification for Latent Class Analysis

Description

This function writes Mplus input files for conducting latent class analysis (LCA) for continuous, count, ordered categorical, and unordered categorical variables. LCA with continuous indicator variables are based on six different variance-covariance structures, while LCA for all other variable types assume local independence. By default, the function conducts LCA with continuous variables

and creates folders in the current working directory for each of the six sets of analysis, writes Mplus input files for conducting LCA with $k = 1$ to $k = 6$ classes into these folders, and writes the matrix or data frame specified in `x` into a Mplus data file in the current working directory. Optionally, all models can be estimated by setting the argument `mplus.run` to `TRUE`.

Usage

```
mplus.lca(x, ind = NULL,
  type = c("continuous", "count", "categorical", "nominal"), cluster = NULL,
  folder = c("A_Invariant-Theta_Diagonal-Sigma",
    "B_Varying-Theta_Diagonal-Sigma",
    "C_Invariant-Theta_Invariant-Unrestricted-Sigma",
    "D_Invariant-Theta_Varying-Unrestricted-Sigma",
    "E_Varying-Theta_Invariant-Unrestricted-Sigma",
    "F_Varying-Theta_Varying-Unrestricted-Sigma"),
  file = "Data_LCA.dat", write = c("all", "folder", "data", "input"),
  useobservations = NULL, missing = -99, classes = 6, estimator = "MLR",
  starts = c(100, 50), stiterations = 10, lrtbootstrap = 1000,
  lrtstarts = c(0, 0, 100, 50), processors = c(8, 8),
  output = c("all", "SVALUES", "CINTERVAL", "TECH7", "TECH8", "TECH11", "TECH14"),
  replace.inp = FALSE, mplus.run = FALSE, Mplus = "Mplus",
  replace.out = c("always", "never", "modified"), check = TRUE)
```

Arguments

<code>x</code>	a matrix or data frame. Note that all variable names must be no longer than 8 character.
<code>ind</code>	a character vector indicating the variables names of the latent class indicators in <code>x</code> .
<code>type</code>	a character string indicating the variable type of the latent class indicators, i.e., "continuous" (default) for continuous variables, "count" for count variables, "categorical" for binary or ordered categorical variables, and "nominal" for unordered categorical variables. Note that it is not possible to mix different variable types in the analysis.
<code>cluster</code>	a character string indicating the cluster variable in the matrix or data frame specified in <code>x</code> representing the nested grouping structure for computing cluster-robust standard errors. Note that specifying a cluster variables does not have any effect on the information criteria, but on the Vuong-Lo-Mendell-Rubin likelihood ratio test of model fit.
<code>folder</code>	a character vector with six character strings for specifying the names of the six folder representing different variance-covariance structures for conducting LCA with continuous indicator variables. There is only one folder for LCA with all other variable types which is called "LCA_1-x_Classes" with <code>x</code> being the maximum number of classes specified in the argument <code>classes</code> .
<code>file</code>	a character string naming the Mplus data file with or without the file extension '.dat', e.g., "Data_LCA.dat" (default) or "Data_LCA".

write	a character string or character vector indicating whether to create the six folders specified in the argument <code>folder</code> ("folder"), to write the matrix or data frame specified in <code>x</code> into a Mplus data file ("data"), and write the Mplus input files into the six folders specified in the argument <code>folder</code> ("input"). By default, the function creates the folders, writes the Mplus data file, and writes the Mplus input files into the folders.
useobservations	a character string indicating the conditional statement to select observations.
missing	a numeric value or character string representing missing values (NA) in the Mplus data set. This values or character string will be specified in the Mplus input file as <code>MISSING IS ALL(missing)</code> . By default, -99 is used to represent missing values.
classes	an integer value specifying the maximum number of classes for the latent class analysis. By default, LCA with a maximum of 6 classes is specified (i.e., $k = 1$ to $k = 6$).
estimator	a character string for specifying the <code>ESTIMATOR</code> option in Mplus. By default, the estimator "MLR" is used.
starts	a vector with two integer values for specifying the <code>STARTS</code> option in Mplus. The first number represents the number of random sets of starting values to generate in the initial stage and the second number represents the optimizations to use in the final stage. By default, 500 random sets of starting values are generated and 100 optimizations are carried out in the final stage.
stiterations	an integer value specifying the <code>STITERATIONS</code> option in Mplus. The numeric value represents the maximum number of iterations allowed in the initial stage. By default, 50 iterations are requested.
lrtbootstrap	an integer value for specifying the <code>LRTBOOTSTRAP</code> option in Mplus when requesting a parametric bootstrapped likelihood ratio test (i.e., <code>output = "TECH14"</code>). The value represents the number of bootstrap draws to be used in estimating the p -value of the parametric bootstrapped likelihood ratio test. By default, 1000 bootstrap draws are requested.
lrtstarts	a vector with four integer values for specifying the <code>LRTSTARTS</code> option in Mplus when requesting a parametric bootstrapped likelihood ratio test (i.e., <code>output = "TECH14"</code>). The values specify the number of starting values to use in the initial stage and the number of optimizations to use in the final stage for the $k - 1$ and k classes model when the data generated by bootstrap draws are analyzed. By default, 0 random sets of starting values in the initial stage and 0 optimizations in the final stage are used for the $k - 1$ classes model and 100 random sets of starting values in the initial stage and 50 optimizations in the final stage are used for the k class model.
processors	a vector of one or two integer values for specifying the <code>PROCESSORS</code> option in Mplus. The values specifies the number of processors and threads to be used for parallel computing to increase computational speed. By default, 8 processors and threads are used for parallel computing.
output	a character string or character vector specifying the <code>TECH</code> options in the <code>OUTPUT</code> section in Mplus, i.e., <code>SVALUES</code> to request input statements that contain parameter estimates from the analysis, <code>CINTERVAL</code> to request confidence intervals,

	TECH7 to request sample statistics for each class using raw data weighted by the estimated posterior probabilities for each class, TECH8 to request the optimization history in estimating the model, TECH11 to request the Lo-Mendell-Rubin likelihood ratio test of model fit, and TECH14 to request a parametric bootstrapped likelihood ratio test. By default, SVALUES and TECH11 are requested. Note that TECH11 is only available for the MLR estimator.
replace.inp	logical: if TRUE, all existing input files in the folder specified in the argument folder are replaced.
mplus.run	logical: if TRUE, all models in the folders specified in the argument folder are estimated by using the <code>mplus.run</code> function in the R package <code>misty</code> .
Mplus	a character string for specifying the name or path of the Mplus executable to be used for running models. This covers situations where Mplus is not in the system's path, or where one wants to test different versions of the Mplus program. Note that there is no need to specify this argument for most users since it has intelligent defaults.
replace.out	a character string for specifying three settings, i.e., "always" to run all models regardless of whether an output file for the model exists, "never" to not run any model that has an existing output file, and "modified" (default) to only runs a model if the modified date for the input file is more recent than the output file modified date.
check	logical: if TRUE (default), argument specification is checked.

Details

Latent class analysis (LCA) is a model-based clustering and classification method used to identify qualitatively different classes of observations which are unknown and must be inferred from the data. LCA can accommodate continuous, count, binary, ordered categorical, and unordered categorical indicators. LCA with continuous indicator variables are also known as latent profile analysis (LPA). In LPA, the within-profile variance-covariance structures represent different assumptions regarding the variance and covariance of the indicator variables both within and between latent profiles. As the best within-profile variance-covariance structure is not known a priori, all of the different structures must be investigated to identify the best model (Masyn, 2013). This function specifies six different variance-covariance structures labeled A to F (see Table 1 in Patterer et al, 2023):

Model A The within-profile variance is constrained to be profile-invariant and covariances are constrained to be 0 in all profiles (i.e., equal variances across profiles and no covariances among indicator variables). This is the default setting in Mplus.

Model B The within-profile variance is profile-varying and covariances are constrained to be 0 in all profiles (i.e., unequal variances across profiles and no covariances among indicator variables).

Model C The within-profile variance is constrained to be profile-invariant and covariances are constrained to be equal in all profiles (i.e., equal variances and covariances across profiles).

Model D The within-profile variance is constrained to be profile-invariant and covariances are profile-varying (i.e., equal variances across profiles and unequal covariances across profiles).

Model E The within-profile variances are profile-varying and covariances are constrained to be equal in all profiles (i.e., unequal variances across profiles and equal covariances across profiles).

Model F The within-class variance and covariances are both profile-varying (i.e., unequal variances and covariances across profiles).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>x</code>	matrix or data frame specified in the argument <code>x</code>
<code>args</code>	specification of function arguments
<code>result</code>	list with six entries for each of the variance-covariance structures and Mplus inputs based on different number of profiles in case of continuous indicators or list of Mplus inputs based on different number of classes in case of count, ordered or unordered categorical indicators.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Masyn, K. E. (2013). Latent class analysis and finite mixture modeling. In T. D. Little (Ed.), *The Oxford handbook of quantitative methods: Statistical analysis* (pp. 551–611). Oxford University Press.
- Muthen, L. K., & Muthen, B. O. (1998-2017). *Mplus User's Guide* (8th ed.). Muthen & Muthen.
- Patterer, A. S., Yanagida, T., Kühnel, J., & Korunka, C. (2023). Daily receiving and providing of social support at work: Identifying support exchange patterns in hierarchical data. *Journal of Work and Organizational Psychology*, 32(4), 489-505. <https://doi.org/10.1080/1359432X.2023.2177537>

See Also

`read.mplus`, `write.mplus`, `mplus`, `mplus.update`, `mplus.print`, `mplus.plot`, `mplus.bayes`, `mplus.run`

Examples

```
## Not run:

# Load data set "HolzingerSwineford1939" in the lavaan package
data("HolzingerSwineford1939", package = "lavaan")

#-----
# Example 1: LCA with k = 1 to k = 8 profiles, continuous indicators
# Input statements that contain parameter estimates
```

```
# Vuong-Lo-Mendell-Rubin LRT and bootstrapped LRT
mplus.lca(HolzingerSwineford1939, ind = c("x1", "x2", "x3", "x4"),
          classes = 8, output = c("SVALUES", "TECH11", "TECH14"))

#-----
# Example 22: LCA with k = 1 to k = 6 profiles, ordered categorical indicators
# Select observations with ageyr <= 13
# Estimate all models in Mplus
mplus.lca(round(HolzingerSwineford1939[, -5]), ind = c("x1", "x2", "x3", "x4"),
          type = "categorical", useobservations = "ageyr <= 13",
          mplus.run = TRUE)

## End(Not run)
```

mplus.plot

Plot Mplus GH5 File

Description

This function uses the `h5file` function in the **hdf5r** package to read a Mplus GH5 file that is requested by the command `PLOT: TYPE IS PLOT2` in Mplus to display trace plots, posterior distribution plots, autocorrelation plots, posterior predictive check plots based on the "bayesian_data" section, and the loop plot based on the "loop_data" section of the Mplus GH5 file. By default, the function displays trace plots if the "bayesian_data" section is available in the Mplus GH5 File. Otherwise, the function plots the loop plot if the "loop_data" section is available in the Mplus GH5 file.

Usage

```
mplus.plot(x, plot = c("none", "trace", "post", "auto", "ppc", "loop"),
           param = c("all", "on", "by", "with", "inter", "var", "r2", "new"),
           std = c("all", "none", "stdyx", "stdy", "std"), burnin = TRUE,
           point = c("all", "none", "m", "med", "map"),
           ci = c("none", "eti", "hdi"), chain = 1, conf.level = 0.95,
           hist = TRUE, density = TRUE, area = TRUE, alpha = 0.4,
           fill = "gray85", facet.nrow = NULL, facet.ncol = NULL,
           facet.scales = c("fixed", "free", "free_x", "free_y"),
           xlab = NULL, ylab = NULL, xlim = NULL, ylim = NULL,
           xbreaks = ggplot2::waiver(), ybreaks = ggplot2::waiver(),
           xexpand = ggplot2::waiver(), yexpand = ggplot2::waiver(),
           palette = "Set 2", binwidth = NULL, bins = NULL,
           density.col = "#0072B2", shape = 21,
           point.col = c("#CC79A7", "#D55E00", "#009E73"),
           linewidth = 0.6, linetype = "dashed", line.col = "black",
           bar.col = "black", bar.width = 0.8, plot.margin = NULL,
           legend.title.size = 10, legend.text.size = 10, legend.box.margin = NULL,
           saveplot = c("all", "none", "trace", "post", "auto", "ppc", "loop"),
           filename = "Mplus_Plot.pdf",
```

```
file.plot = c("_TRACE", "_POST", "_AUTO", "_PPC", "_LOOP"),
width = NA, height = NA, units = c("in", "cm", "mm", "px"),
dpi = 600, check = TRUE)
```

Arguments

<code>x</code>	a character string indicating the name of the Mplus GH5 file (HDF5 format) with or without the file extension .gh5, e.g., "Mplus_Plot.gh5" or "Mplus_Plot". Alternatively, a <code>misty</code> object of type <code>mplus</code> can be specified, i.e., result object of the <code>mplus.plot()</code> function.
<code>plot</code>	a character string indicating the type of plot to display, i.e., "none" for not displaying any plot, "trace" (default) for displaying trace plots, "post" for displaying posterior distribution plots, "auto" for displaying autocorrelation plots, "ppc" for displaying posterior predictive check plots, and "loop" for displaying the loop plot. The default setting is "trace" if the "bayesian_data" section is available in the Mplus GH5 file. Otherwise, the default setting switches to "loop".
<code>param</code>	character vector indicating which parameters to print for the trace plots, posterior distribution plots, and autocorrelation plots, i.e., "all" for all parameters, "on" (default), for regression slopes, "by" for factor loadings, "with" for covariances, "inter" for intercepts and thresholds, "var" for (residual) variances, "r2" for r-square, and "new" for parameters not in the analysis model specified in the NEW option. The default setting is "on" if regression slopes are available. Otherwise, the default setting switches to "by" and to "with" if factor loadings are not available.
<code>std</code>	a character vector indicating the standardized parameters to print for the trace plots, posterior distribution plots, and autocorrelation plots, i.e., "all" for all standardized parameters, "none" (default) for not printing any standardized parameters, "stdyx" for StdYX standardized parameters, "stdy" for StdY standardized parameters, and "std" for StdX standardized parameters.
<code>burnin</code>	logical: if FALSE, the first half of each chain is discarded as being part of the burnin phase when displaying trace plots. The default setting for <code>plot = "trace"</code> is TRUE. Note that the first half of each chain is always discarded when displaying posterior distribution plots (<code>plot = "post"</code>) regardless of the setting of the argument <code>burnin</code> .
<code>point</code>	a character vector indicating the point estimate(s) to be displayed in the posterior distribution plots, i.e., "all" for all point estimates, "none" for not displaying any point estimates, "m" for the posterior mean estimate, "med" (default) for the posterior median estimate, and "map" for the maximum a posteriori estimate.
<code>ci</code>	a character string indicating the type of credible interval to be displayed in the posterior distribution plots, i.e., "none" for not displaying any credible intervals, "eti" (default) for displaying the equal-tailed intervals and "hdi" for displaying the highest density interval.
<code>chain</code>	a numerical value indicating the chain to be used for the autocorrelation plots. By default, the first chain is used.
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the credible interval (default is 0.95).

hist	logical: if TRUE (default), histograms are drawn in the posterior probability plots.
density	logical: if TRUE (default), density curves are drawn in the posterior probability plots.
area	logical: if TRUE (default), statistical not significant and statistical significant area is filled with a different color and vertical lines are drawn.
alpha	a numeric value between 0 and 1 for the alpha argument (default is 0.4) for the annotate, geom_histogram, geom_bar, and geom_ribbon function.
fill	a character string indicating the color for the "fill" argument (default is "gray85") for the annotate, geom_histogram, geom_bar, and geom_point functions.
facet.nrow	a numeric value indicating the nrow argument (default is NULL) for the facet_wrap function.
facet.ncol	a numeric value indicating the ncol argument (default is 2) for the facet_wrap function.
facet.scales	a character string indicating the scales argument (default is "free") for the facet_wrap function.
xlab	a character string indicating the name argument for the scale_x_continuous function. Note that the default setting depends on the type of plot, e.g., "" for the trace plots and "Lag" for the autocorrelation plots.
ylab	a character string indicating the name argument for the scale_y_continuous function. Note that the default setting depends on the type of plot, e.g., "" for the trace plots and "Autocorrelation" for the autocorrelation plots.
xlim	a numeric vector with two elements indicating the limits argument (default is NULL) for the scale_x_continuous function.
ylim	a numeric vector with two elements indicating the limits argument (default is NULL) for the scale_y_continuous function.
xbreaks	a numeric vector indicating the breaks argument (default is ggplot2::waiver()) for the scale_x_continuous function.
ybreaks	a numeric vector indicating the breaks argument (default is ggplot2::waiver()) for the scale_y_continuous function.
xexpand	a numeric vector with two elements indicating the expand argument (default is (0.02, 0)) for the scale_x_continuous function.
yexpand	a numeric vector with two elements indicating the expand argument for the scale_y_continuous function. Note that the default setting depends on the type of plot, e.g., (0.02, 0) for the trace plots and expansion(mult = c(0, 0.05)) for the posterior distribution plots.
palette	a character string indicating the palette name (default is "Set 2") for the hcl.colors function. Note that the character string must be one of hcl.pals().
binwidth	a numeric value indicating the binwidth argument (default is to use the number of bins in bins argument) for the geom_histogram function.
bins	a numeric value indicating the bins argument (default is 30) for the geom_histogram function.
density.col	a character string indicating the color argument (default is "#0072B2") for the geom_density function.

shape	a numeric value indicating the shape argument (default is 21) for the <code>geom_point</code> function.
point.col	a character vector with three elements indicating the values argument (default is <code>c("#CC79A7", "#D55E00", "#009E73")</code>) for the <code>scale_color_manual</code> function.
linewidth	a numeric value indicating the linewidth argument (default is 0.6) for the <code>geom_vline</code> function.
linetype	a numeric value indicating the linetype argument (default is "dashed") for the <code>geom_vline</code> function.
line.col	a character string indicating the color argument (default is "black") for the <code>geom_vline</code> function.
bar.col	a character string indicating the color argument (default is "black") for the <code>geom_bar</code> function.
bar.width	a character string indicating the width argument (default = 0.8) for the <code>geom_bar</code> function.
plot.margin	a numeric vector indicating the <code>plot.margin</code> argument for the <code>theme</code> function. Note that the default setting depends on the type of the plot, e.g., <code>c(4, 15, -10, 0)</code> for the trace plots, and <code>c(4, 15, 4, 4)</code> for the autocorrelation plots.
legend.title.size	a numeric value indicating the <code>legend.title</code> argument (default is <code>element_text(size = 10)</code>) for the <code>theme</code> function.
legend.text.size	a numeric value indicating the <code>legend.text</code> argument (default is <code>element_text(size = 10)</code>) for the <code>theme</code> function.
legend.box.margin	a numeric vector indicating the <code>legend.box.margin</code> argument for the <code>theme</code> function. Note that the default setting depends on the type of plot, e.g., <code>c(-16, 6, 6, 6)</code> for the trace plots, and <code>c(-25, 6, 6, 6)</code> for the posterior distribution plots with displaying point estimates.
saveplot	a character vector indicating the plot to be saved, i.e., "all" for saving all plots, "none" (default) for not saving any plots, "trace" for saving the trace plots, "post" for the saving the posterior distribution plots, "auto" for saving the autocorrelation plots, "ppc" for saving the posterior predictive check plots, and "loop" for saving the loop plot.
filename	a character string indicating the filename argument (default is "Mplus_Plot.pdf") including the file extension for the <code>ggsave</code> function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the filename argument.
file.plot	a character vector with five elements for distinguishing different types of plots. By default, the character string specified in the argument "filename" ("Mplus_Plot") is concatenated with "_TRACE" ("Mplus_Plot_TRACE") for the trace plots, "_POST" ("Mplus_Plot_POST") for the posterior distribution plots, "_AUTO" ("Mplus_Plot_AUTO") for the autocorrelation plots, "_PPC" ("Mplus_Plot_PPC") for the posterior predictive check plots, and "_LOOP" ("Mplus_Plot_LOOP") for the loop plot.

width	a numeric value indicating the width argument (default is the size of the current graphics device) for the ggsave function.
height	a numeric value indicating the height argument (default is the size of the current graphics device) for the ggsave function.
units	a character string indicating the units argument (default is in) for the ggsave function.
dpi	a numeric value indicating the dpi argument (default is 600) for the ggsave function.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
x	Mplus GH5 file
args	specification of function arguments
data	list with posterior distribution of each parameter estimate in wide and long format (post), autocorrelation for each parameter estimate in wide and long format (auto), data for the posterior predictive check (ppc), and data for the loop plot (loop)
plot	list with the trace plots (trace, posterior distribution plots (post), autocorrelation plots (auto), posterior predictive check plots (ppc), and loop plot (loop)

Author(s)

Takuya Yanagida

References

Muthen, L. K., & Muthen, B. O. (1998-2017). *Mplus User's Guide* (8th ed.). Muthen & Muthen.

See Also

[read.mplus](#), [write.mplus](#), [mplus](#), [mplus.update](#), [mplus.print](#), [mplus.bayes](#), [mplus.run](#), [mplus.lca](#)

Examples

```
## Not run:

#-----
# Mplus Example 3.18: Moderated Mediation with a Plot of the Indirect Effect

#.....
# Trace Plots

# Example 1a: Default setting
```

```

mplus.plot("ex3.18.gh5")

# Example 1b: Exclude first half of each chain
mplus.plot("ex3.18.gh5", burnin = FALSE)

# Example 1c: Print all parameters
mplus.plot("ex3.18.gh5", param = "all")

# Example 1d: Print user-specified parameters
mplus.plot("ex3.18.gh5", param = "param")

# Example 1e: Arrange panels in three columns
mplus.plot("ex3.18.gh5", ncol = 3)

# Example 1f: Specify "Pastel 1" palette for the hcl.colors function
mplus.plot("ex3.18.gh5", palette = "Pastel 1")

#.....
# Posterior Distribution Plots

# Example 2a: Default setting, i.e., posterior median and equal-tailed interval
mplus.plot("ex3.18.gh5", plot = "post")

# Example 2b: Display posterior mean and maximum a posteriori
mplus.plot("ex3.18.gh5", plot = "post", point = c("m", "map"))

# Example 2c: Display maximum a posteriori and highest density interval
mplus.plot("ex3.18.gh5", plot = "post", point = "map", ci = "hdi")

# Example 2d: Do not display any point estimates and credible interval
mplus.plot("ex3.18.gh5", plot = "post", point = "none", ci = "none")

# Example 2d: Do not display histograms
mplus.plot("ex3.18.gh5", plot = "post", hist = FALSE)

#.....
# Autocorrelation Plots

# Example 3a: Default setting, i.e., first chain
mplus.plot("ex3.18.gh5", plot = "auto")

# Example 3b: Use second chain
mplus.plot("ex3.18.gh5", plot = "auto", chain = 2)

# Example 3b: Modify limits and breaks of the y-axis
mplus.plot("ex3.18.gh5", plot = "auto",
          ylim = c(-0.05, 0.05), ybreaks = seq(-0.1, 0.1, by = 0.025))

#.....
# Posterior Predictive Check Plots

# Example 4a: Default setting, i.e., 95% Interval
mplus.plot("ex3.18.gh5", plot = "ppc")

```

```

# Example 4b: Default setting, i.e., 99% Interval
mplus.plot("ex3.18.gh5", plot = "ppc", conf.level = 0.99)

#.....
# Loop Plot

# Example 5a: Default setting
mplus.plot("ex3.18.gh5", plot = "loop")

# Example 5b: Do not fill area and draw vertical lines
mplus.plot("ex3.18.gh5", plot = "loop", area = FALSE)

#.....
# Save Plots

# Example 6a: Save all plots in pdf format
mplus.plot("ex3.18.gh5", saveplot = "all")

# Example 6b: Save all plots in png format with 300 dpi
mplus.plot("ex3.18.gh5", saveplot = "all", filename = "Mplus_Plot.png", dpi = 300)

# Example 6a: Save loop plot, specify width and height of the plot
mplus.plot("ex3.18.gh5", plot = "none", saveplot = "loop",
           width = 7.5, height = 7)

#-----
# Plot from misty.object

# Create misty.object
object <- mplus.plot("ex3.18.gh5", plot = "none")

# Trace plot
mplus.plot(object, plot = "trace")

# Posterior distribution plot
mplus.plot(object, plot = "post")

# Autocorrelation plot
mplus.plot(object, plot = "auto")

# Posterior predictive check plot
mplus.plot(object, plot = "ppc")

# Loop plot
mplus.plot(object, plot = "loop")

#-----
# Create Plots Manually

# Load ggplot2 package
library(ggplot2)

```

```

# Create misty object
object <- mplus.plot("ex3.18.gh5", plot = "none")

#.....
# Example 7: Trace Plots

# Extract data in long format
data.post <- object$data$post$long

# Extract ON parameters
data.trace <- data.post[grep(" ON ", data.post$param), ]

# Plot
ggplot(data.trace, aes(x = iter, y = value, color = chain)) +
  annotate("rect", xmin = 0, xmax = 15000, ymin = -Inf, ymax = Inf,
    alpha = 0.4, fill = "gray85") +
  geom_line() +
  facet_wrap(~ param, ncol = 2, scales = "free") +
  scale_x_continuous(name = "", expand = c(0.02, 0)) +
  scale_y_continuous(name = "", expand = c(0.02, 0)) +
  scale_colour_manual(name = "Chain",
    values = hcl.colors(n = 2, palette = "Set 2")) +
  theme_bw() +
  guides(color = guide_legend(nrow = 1, byrow = TRUE)) +
  theme(plot.margin = margin(c(4, 15, -10, 0)),
    legend.position = "bottom",
    legend.title = element_text(size = 10),
    legend.text = element_text(size = 10),
    legend.box.margin = margin(c(-16, 6, 6, 6)),
    legend.background = element_rect(fill = "transparent"))

#.....
# Example 8: Posterior Distribution Plots

# Extract data in long format
data.post <- object$data$post$long

# Extract ON parameters
data.post <- data.post[grep(" ON ", data.post$param), ]

# Discard burn-in iterations
data.post <- data.post[data.post$iter > 15000, ]

# Drop factor levels
data.post$param <- droplevels(data.post$param,
  exclude = c("[Y]", "[M]", "Y", "M", "INDIRECT", "MOD"))

# Plot
ggplot(data.post, aes(x = value)) +
  geom_histogram(aes(y = after_stat(density)), color = "black", alpha = 0.4,
    fill = "gray85") +
  geom_density(color = "#0072B2") +
  geom_vline(data = data.frame(param = unique(data.post$param)),

```

```

      stat = tapply(data.post$value, data.post$param, median)),
    aes(xintercept = stat, color = "Median"), linewidth = 0.6) +
  geom_vline(data = data.frame(param = unique(data.post$param),
    low = tapply(data.post$value, data.post$param,
      function(y) quantile(y, probs = 0.025))),
    aes(xintercept = low), linetype = "dashed", linewidth = 0.6) +
  geom_vline(data = data.frame(param = unique(data.post$param),
    upp = tapply(data.post$value, data.post$param,
      function(y) quantile(y, probs = 0.975))),
    aes(xintercept = upp), linetype = "dashed", linewidth = 0.6) +
  facet_wrap(~ param, ncol = 2, scales = "free") +
  scale_x_continuous(name = "", expand = c(0.02, 0)) +
  scale_y_continuous(name = "Probability Density, f(x)",
    expand = expansion(mult = c(0L, 0.05))) +
  scale_color_manual(name = "Point Estimate", values = c(Median = "#D55E00")) +
  labs(caption = "95% Equal-Tailed Interval") +
  theme_bw() +
  theme(plot.margin = margin(c(4, 15, -8, 4)),
    plot.caption = element_text(hjust = 0.5, vjust = 7),
    legend.position = "bottom",
    legend.title = element_text(size = 10),
    legend.text = element_text(size = 10),
    legend.box.margin = margin(c(-30, 6, 6, 6)),
    legend.background = element_rect(fill = "transparent"))

#.....
# Example 9: Autocorrelation Plots

# Extract data in long format
data.auto <- object$data$auto$long

# Select first chain
data.auto <- data.auto[data.auto$chain == 1, ]

# Extract ON parameters
data.auto <- data.auto[grepl(" ON ", data.auto$param), ]

# Plot
ggplot(data.auto, aes(x = lag, y = cor)) +
  geom_bar(stat = "identity", alpha = 0.4, color = "black", fill = "gray85",
    width = 0.8) +
  facet_wrap(~ param, ncol = 2) +
  scale_x_continuous(name = "Lag", breaks = seq(1, 30, by = 2), expand = c(0.02, 0)) +
  scale_y_continuous(name = "Autocorrelation", limits = c(-0.1, 0.1),
    breaks = seq(-0.1, 1., by = 0.05), expand = c(0.02, 0)) +
  theme_bw() +
  theme(plot.margin = margin(c(4, 15, 4, 4)))

#.....
# Example 10: Posterior Predictive Check (PPC) Plots

# Extract data
data.ppc <- object$data$ppc

```

```

# Scatter plot
ppc.scatter <- ggplot(data.ppc, aes(x = obs, y = rep)) +
  geom_point(shape = 21, fill = "gray85") +
  geom_abline(slope = 1) +
  scale_x_continuous("Observed", limits = c(0, 45), breaks = seq(0, 45, by = 5),
    expand = c(0.02, 0)) +
  scale_y_continuous("Recpliated", limits = c(0, 45), breaks = seq(0, 45, by = 5),
    expand = c(0.02, 0)) +
  theme_bw() +
  theme(plot.margin = margin(c(2, 15, 4, 4)))

# Histogram
ppc.hist <- ggplot(data.ppc, aes(x = diff)) +
  geom_histogram(color = "black", alpha = 0.4, fill = "gray85") +
  geom_vline(xintercept = mean(data.ppc$diff), color = "#CC79A7") +
  geom_vline(xintercept = quantile(data.ppc$diff, probs = 0.025),
    linetype = "dashed", color = "#CC79A7") +
  geom_vline(xintercept = quantile(data.ppc$diff, probs = 0.975),
    linetype = "dashed", color = "#CC79A7") +
  scale_x_continuous("Observed - Replicated", expand = c(0.02, 0)) +
  scale_y_continuous("Count", expand = expansion(mult = c(0L, 0.05))) +
  theme_bw() +
  theme(plot.margin = margin(c(2, 15, 4, 4)))

# Combine plots using the patchwork package
patchwork::wrap_plots(ppc.scatter, ppc.hist)

#.....
# Example 11: Loop Plot

# Extract data
data.loop <- object$data$loop

# Plot
plot.loop <- ggplot(data.loop, aes(x = xval, y = estimate)) +
  geom_line(linewidth = 0.6, show.legend = FALSE) +
  geom_line(aes(xval, low)) +
  geom_line(aes(xval, upp)) +
  scale_x_continuous("MOD", expand = c(0.02, 0)) +
  scale_y_continuous("INDIRECT", expand = c(0.02, 0)) +
  scale_fill_manual("Statistical Significance",
    values = hcl.colors(n = 2, palette = "Set 2")) +
  theme_bw() +
  theme(plot.margin = margin(c(4, 15, -6, 4)),
    legend.position = "bottom",
    legend.title = element_text(size = 10),
    legend.text = element_text(size = 10),
    legend.box.margin = margin(-10, 6, 6, 6),
    legend.background = element_rect(fill = "transparent"))

# Significance area
for (i in unique(data.loop$group)) {

```

```

plot.loop <- plot.loop + geom_ribbon(data = data.loop[data.loop$group == i, ],
                                   aes(ymin = low, ymax = upp, fill = sig), alpha = 0.4)

}

# Vertical lines
plot.loop + geom_vline(data = data.loop[data.loop$change == 1, ],
                       aes(xintercept = xval, color = sig), linewidth = 0.6,
                           linetype = "dashed", show.legend = FALSE)

## End(Not run)

```

mplus.print

Print Mplus Output

Description

This function prints the input command sections and the result sections of a Mplus output file (.out) on the R console. By default, the function prints selected result sections, e.g., short Summary of Analysis, short Summary of Data, Model Fit Information, and Model Results.

Usage

```

mplus.print(x, print = c("all", "input", "result"),
            input = c("all", "default", "data", "variable", "define",
                      "analysis", "model", "montecarlo", "mod.pop", "mod.cov",
                      "mod.miss", "message"),
            result = c("all", "default", "summary.analysis.short",
                      "summary.data.short", "random.starts", "summary.fit",
                      "mod.est", "fit", "class.count", "classif",
                      "mod.result", "total.indirect"),
            exclude = NULL, variable = FALSE, not.input = TRUE, not.result = TRUE,
            write = NULL, append = TRUE, check = TRUE, output = TRUE)

```

Arguments

<code>x</code>	a character string indicating the name of the Mplus output file with or without the file extension .out, e.g., "Mplus_Output.out" or "Mplus_Output". Alternatively, a misty.object of type mplus can be specified, i.e., result object of the mplus.print(), mplus() or mplus.update() function.
<code>print</code>	a character vector indicating which section to show, i.e. "all" for input and result sections, "input" for input command section only, and "result" (default) for result sections only
<code>input</code>	a character vector specifying Mplus input command sections
<code>result</code>	a character vector specifying Mplus result sections included in the output (see 'Details').

exclude	a character vector specifying Mplus input command or result sections excluded from the output (see 'Details').
variable	logical: if TRUE, names of the variables in the data set (NAMES option) specified in the VARIABLE: command section are shown. By default, names of the variables in the data set are excluded from the output unless all variables are used in the analysis (i.e., no USEVARIABLES option specified in the Mplus input file).
not.input	logical: if TRUE (default), character vector indicating the input commands not requested are shown on the console.
not.result	logical: if TRUE (default), character vector indicating the result sections not requested are shown on the console.
write	a character string naming a file for writing the output into a text file with file extension ".txt" (e.g., "Output.txt").
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

Input Command Sections Following input command sections can be selected by using the input argument or excluded by using the exclude argument:

- "title" for the TITLE command used to provide a title for the analysis.
- "data" for the DATA command used to provide information about the data set to be analyzed.
- "data.imp" for the DATA IMPUTATION command used to create a set of imputed data sets using multiple imputation methodology.
- "data.wl" for the DATA WIDETOLONG command used to rearrange data from a multivariate wide format to a univariate long format.
- "data.lw" for the DATA LONGTOWIDE command used to rearrange a univariate long format to a multivariate wide format.
- "data.tp" for the DATA TWOPART command used to create a binary and a continuous variable from a continuous variable with a floor effect for use in two-part modeling.
- "data.miss" for the DATA MISSING command used to create a set of binary variables that are indicators of missing data or dropout for another set of variables.
- "data.surv" for the DATA SURVIVAL command used to create variables for discrete-time survival modeling.
- "data.coh" for the DATA COHORT command used to rearrange longitudinal data from a format where time points represent measurement occasions to a format where time points represent age or another time-related variable.
- "variable" for the VARIABLE command used to provide information about the variables in the data set to be analyzed.
- "define" for the DEFINE command used to transform existing variables and to create new variables.
- "analysis" for the ANALYSIS command used to describe the technical details for the analysis.

- "model" MODEL for the command used to describe the model to be estimated.
- "mod.ind" for the MODEL INDIRECT command used to request indirect and directed effects and their standard errors.
- "mod.test" for the MODEL TEST command used to test restrictions on the parameters in the MODEL and MODEL CONSTRAINT commands using the Wald chi-square test.
- "mod.prior" for the MODEL PRIORS command used with ESTIMATOR IS BAYES to specify the prior distribution for each parameter.
- "montecarlo" for the MONTECARLO command used to set up and carry out a Monte Carlo simulation study.
- "mod.pop" for the MODEL POPULATION command used to provide the population parameter values to be used in data generation using the options of the MODEL command.
- "mod.cov" for the MODEL COVERAGE used to provide the population parameter values to be used for computing coverage.
- "mod.miss" for the MODEL MISSING command used to provide information about the population parameter values for the missing data model to be used in the generation of data.
- "output" for the for the OUTPUT command used to request additional output beyond that included as the default.
- "savedata" for the SAVEDATA command used to save the analysis data and/or a variety of model results in an ASCII file for future use.
- "plot" for the PLOT command used to requested graphical displays of observed data and analysis results.
- "message" for warning and error messages that have been generated by the program after the input command sections.

Note that all input command sections are requested by specifying `input = "all"`. The input argument is also used to select one (e.g., `input = "model"`) or more than one input command sections (e.g., `input = c("analysis", "model")`), or to request input command sections in addition to the default setting (e.g., `input = c("default", "output")`). The exclude argument is used to exclude input command sections from the output (e.g., `exclude = "variable"`).

Result Sections Following result sections can be selected by using the result argument or excluded by using the exclude argument:

- "summary.analysis" for the SUMMARY OF ANALYSIS section..
- "summary.analysis.short" for a short SUMMARY OF ANALYSIS section including the number of observations, number of groups, estimator, and optimization algorithm.
- "summary.data" for the SUMMARY OF DATA section indicating.
- "summary.data.short" for a short SUMMARY OF DATA section including number of clusters, average cluster size, and estimated intraclass correlations.
- "prop.count" for the UNIVARIATE PROPORTIONS AND COUNTS FOR CATEGORICAL VARIABLES section.
- "summary.censor" for the SUMMARY OF CENSORED LIMITS section.
- "prop.zero" for the COUNT PROPORTION OF ZERO, MINIMUM AND MAXIMUM VALUES section.
- "crosstab" for the CROSSTABS FOR CATEGORICAL VARIABLES section.
- "summary.miss" for the SUMMARY OF MISSING DATA PATTERNS section.
- "coverage" for the COVARIANCE COVERAGE OF DATA section.

- "basic" for the RESULTS FOR BASIC ANALYSIS section.
- "sample.stat" for the SAMPLE STATISTICS section.
- "uni.sample.stat" for the UNIVARIATE SAMPLE STATISTICS section.
- "random.starts" for the RANDOM STARTS RESULTS section.
- "summary.fit" for the SUMMARY OF MODEL FIT INFORMATION section.
- "mod.est" for the THE MODEL ESTIMATION TERMINATED NORMALLY message and warning messages from the model estimation.
- "fit" for the MODEL FIT INFORMATION section.
- "class.count" for the FINAL CLASS COUNTS AND PROPORTIONS FOR THE LATENT CLASSES section.
- "ind.means" for the LATENT CLASS INDICATOR MEANS AND PROBABILITIES section.
- "trans.prob" for the LATENT TRANSITION PROBABILITIES BASED ON THE ESTIMATED MODEL section.
- "classif" for the CLASSIFICATION QUALITY section.
- "mod.result" for the MODEL RESULTS and RESULTS FOR EXPLORATORY FACTOR ANALYSIS section.
- "odds.ratio" for the LOGISTIC REGRESSION ODDS RATIO RESULTS section.
- "prob.scale" for the RESULTS IN PROBABILITY SCALE section.
- "ind.odds.ratio" for the LATENT CLASS INDICATOR ODDS RATIOS FOR THE LATENT CLASSES section.
- "alt.param" for the ALTERNATIVE PARAMETERIZATIONS FOR THE CATEGORICAL LATENT VARIABLE REGRESSION section.
- "irt.param" for the IRT PARAMETERIZATION section.
- "brant.wald" for the BRANT WALD TEST FOR PROPORTIONAL ODDS section.
- "std.mod.result" for the STANDARDIZED MODEL RESULTS section.
- "rsquare" for the R-SQUARE section.
- "total.indirect" for the TOTAL, TOTAL INDIRECT, SPECIFIC INDIRECT, AND DIRECT EFFECTS section.
- "std.total.indirect" for the STANDARDIZED TOTAL, TOTAL INDIRECT, SPECIFIC INDIRECT, AND DIRECT EFFECTS section.
- "std.mod.result.cluster" for the WITHIN-LEVEL STANDARDIZED MODEL RESULTS FOR CLUSTER section.
- "fs.comparison" for the BETWEEN-LEVEL FACTOR SCORE COMPARISONS section.
- "conf.mod.result" for the CONFIDENCE INTERVALS OF MODEL RESULTS section.
- "conf.std.conf" for the CONFIDENCE INTERVALS OF STANDARDIZED MODEL RESULTS section.
- "conf.total.indirect" for the CONFIDENCE INTERVALS OF TOTAL, TOTAL INDIRECT, SPECIFIC INDIRECT, AND DIRECT EFFECTS section.
- "conf.odds.ratio" for the CONFIDENCE INTERVALS FOR THE LOGISTIC REGRESSION ODDS RATIO RESULTS section.
- "modind" for the MODEL MODIFICATION INDICES section.
- "resid" for the RESIDUAL OUTPUT section.
- "logrank" for the LOGRANK OUTPUT section.
- "tech1" for the TECHNICAL 1 OUTPUT section.

- "tech2" for the TECHNICAL 2 OUTPUT section.
- "tech3" for the TECHNICAL 3 OUTPUT section.
- "h1.tech3" for the H1 TECHNICAL 3 OUTPUT section.
- "tech4" for the TECHNICAL 4 OUTPUT section.
- "tech5" for the TECHNICAL 5 OUTPUT section.
- "tech6" for the TECHNICAL 6 OUTPUT section.
- "tech7" for the TECHNICAL 7 OUTPUT section.
- "tech8" for the TECHNICAL 8 OUTPUT section.
- "tech9" for the TECHNICAL 9 OUTPUT section.
- "tech10" for the TECHNICAL 10 OUTPUT section.
- "tech11" for the TECHNICAL 11 OUTPUT section.
- "tech12" for the TECHNICAL 12 OUTPUT section.
- "tech13" for the TECHNICAL 13 OUTPUT section.
- "tech14" for the TECHNICAL 14 OUTPUT section.
- "tech15" for the TECHNICAL 15 OUTPUT section.
- "tech16" for the TECHNICAL 16 OUTPUT section.
- "svalues" for the MODEL COMMAND WITH FINAL ESTIMATES USED AS STARTING VALUES section.
- "stat.fscores" for the SAMPLE STATISTICS FOR ESTIMATED FACTOR SCORES section.
- "summary.fscores" for the SUMMARY OF FACTOR SCORES section.
- "pv" for the SUMMARIES OF PLAUSIBLE VALUES section.
- "plotinfo" for the PLOT INFORMATION section.
- "saveinfo" for the SAVEDATA INFORMATION section.

Note that all result sections are requested by specifying `result = "all"`. The `result` argument is also used to select one (e.g., `result = "mod.result"`) or more than one result sections (e.g., `result = c("mod.result", "std.mod.result")`), or to request result sections in addition to the default setting (e.g., `result = c("default", "odds.ratio")`). The `exclude` argument is used to exclude result sections from the output (e.g., `exclude = "mod.result"`).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>x</code>	character string or <code>misty.object</code>
<code>args</code>	specification of function arguments
<code>print</code>	print objects
<code>notprint</code>	character vectors indicating the input commands and result sections not requested
<code>result</code>	list with input command sections (input) and result sections (input)

Author(s)

Takuya Yanagida

References

Muthen, L. K., & Muthen, B. O. (1998-2017). *Mplus User's Guide* (8th ed.). Muthen & Muthen.

See Also

[read.mplus](#), [write.mplus](#), [mplus](#), [mplus.update](#), [mplus.plot](#), [mplus.bayes](#), [mplus.run](#), [mplus.lca](#)

Examples

```
## Not run:

#-----
# Mplus Example 3.1: Linear Regression

# Example 1a: Default setting
mplus.print("ex3.1.out")

# Example 1b: Print result section only
mplus.print("ex3.1.out", print = "result")

# Example 1c: Print MODEL RESULTS only
mplus.print("ex3.1.out", print = "result", result = "mod.result")

# Example 1d: Print UNIVARIATE SAMPLE STATISTICS in addition to the default setting
mplus.print("ex3.1.out", result = c("default", "uni.sample.stat"))

# Example 1e: Exclude MODEL FIT INFORMATION section
mplus.print("ex3.1.out", exclude = "fit")

# Example 1f: Print all result sections, but exclude MODEL FIT INFORMATION section
mplus.print("ex3.1.out", result = "all", exclude = "fit")

# Example 1g: Print result section in a different order
mplus.print("ex3.1.out", result = c("mod.result", "fit", "summary.analysis"))

#-----
# misty.object of type 'mplus.print'

# Example 2
# Create misty.object
object <- mplus.print("ex3.1.out", output = FALSE)

# Print misty.object
mplus.print(object)

#-----
# Write Results

# # Example 3: Write Results into a text file
mplus.print("ex3.1.out", write = "Output_3-1.txt")

## End(Not run)
```

mplus.run

*Run Mplus Models***Description**

This function runs a group of Mplus models (.inp files) located within a single directory or nested within subdirectories.

Usage

```
mplus.run(target = getwd(), recursive = FALSE, filefilter = NULL, show.out = FALSE,
  replace.out = c("always", "never", "modified"), message = TRUE,
  logFile = NULL, Mplus = .detect.mplus(), killOnFail = TRUE,
  local_tmpdir = FALSE, check = TRUE)
```

Arguments

target	a character string indicating the directory containing Mplus input files (.inp) to run or the single .inp file to be run. May be a full path, relative path, or a filename within the working directory.
recursive	logical: if TRUE, run all models nested in subdirectories within directory. Not relevant if target is a single file.
filefilter	a Perl regular expression (PCRE-compatible) specifying particular input files to be run within directory. See regex or http://www.pcre.org/pcre.txt for details about regular expression syntax. Not relevant if target is a single file.
show.out	logical: if TRUE, estimation output (TECH8) is show on the R console. Note that if run within Rgui, output will display within R, but if run via Rterm, a separate window will appear during estimation.
replace.out	a character string for specifying three settings: "always" (default), which runs all models, regardless of whether an output file for the model exists, "never", which does not run any model that has an existing output file, and "modified", which only runs a model if the modified date for the input file is more recent than the output file modified date.
message	logical: if TRUE, message Running model: and System command: is printed on the console.
logFile	a character string specifying a file that records the settings passed into the function and the models run (or skipped) during the run.
Mplus	a character string for specifying the name or path of the Mplus executable to be used for running models. This covers situations where Mplus is not in the system's path, or where one wants to test different versions of the Mplus program. Note that there is no need to specify this argument for most users since it has intelligent defaults.
killOnFail	logical: if TRUE (default), all processes named mplus.exe when mplus.run() does not terminate normally are killed. Windows only.

local_tmpdir	logical: if TRUE, the TMPDIR environment variable is set to the location of the .inp file prior to execution. This is useful in Monte Carlo studies where many instances of Mplus may run in parallel and we wish to avoid collisions in temporary files among processes. Linux/Mac only.
check	logical: if TRUE (default), argument specification, convergence and model identification is checked.

Value

None.

Note

This function is a copy of the runModels() function in the **MplusAutomation** package by Michael Hallquist and Joshua Wiley (2018).

Author(s)

Michael Hallquist and Joshua Wiley

References

Hallquist, M. N. & Wiley, J. F. (2018). MplusAutomation: An R package for facilitating large-scale latent variable analyses in Mplus. *Structural Equation Modeling: A Multidisciplinary Journal*, 25, 621-638. <https://doi.org/10.1080/10705511.2017.1402334>.

Muthen, L. K., & Muthen, B. O. (1998-2017). *Mplus User's Guide* (8th ed.). Muthen & Muthen.

See Also

`read.mplus`, `write.mplus`, `mplus`, `mplus.update`, `mplus.print`, `mplus.plot`, `mplus.bayes`, `mplus.lca`

Examples

```
## Not run:

# Example 1: Run Mplus models located within a single directory
mplus.run(Mplus = "C:/Program Files/Mplus/Mplus.exe")

# Example 2: Run Mplus models located nested within subdirectories
mplus.run(recursive = TRUE,
          Mplus = "C:/Program Files/Mplus/Mplus.exe")

## End(Not run)
```

mplus.update

*Mplus Input Updating***Description**

This function updates specific input command sections of a `misty.object` of type `mplus` to create an updated Mplus input file, run the updated input file by using the `mplus.run()` function, and print the updated Mplus output file by using the `mplus.print()` function.

Usage

```
mplus.update(x, update, file = "Mplus_Input_Update.inp", comment = FALSE,
             replace.inp = TRUE, mplus.run = TRUE,
             show.out = FALSE, replace.out = c("always", "never", "modified"),
             print = c("all", "input", "result"),
             input = c("all", "default", "data", "variable", "define",
                       "analysis", "model", "montecarlo", "mod.pop", "mod.cov",
                       "mod.miss", "message"),
             result = c("all", "default", "summary.analysis.short",
                       "summary.data.short", "random.starts", "summary.fit",
                       "mod.est", "fit", "class.count", "classif",
                       "mod.result", "total.indirect"),
             exclude = NULL, variable = FALSE, not.input = TRUE, not.result = TRUE,
             write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>x</code>	<code>misty.object</code> object of type <code>mplus</code> .
<code>update</code>	a character string containing the updated input command sections.
<code>file</code>	a character string indicating the name of the updated Mplus input file with or without the file extension <code>.inp</code> , e.g., <code>"Mplus_Input_Update.inp"</code> or <code>"Mplus_Input_Update"</code> .
<code>comment</code>	logical: if <code>FALSE</code> (default), comments (i.e., text after the <code>!</code> symbol) are removed from the input text specified in the argument <code>x</code> .
<code>replace.inp</code>	logical: if <code>TRUE</code> (default), an existing input file will be replaced.
<code>mplus.run</code>	logical: if <code>TRUE</code> , the input file specified in the argument <code>file</code> containing the input text specified in the argument <code>x</code> is run using the <code>mplus.run</code> function.
<code>show.out</code>	logical: if <code>TRUE</code> , estimation output (TECH8) is show on the R console. Note that if run within Rgui, output will display within R, but if run via Rterm, a separate window will appear during estimation.
<code>replace.out</code>	a character string for specifying three settings: <code>"always"</code> (default), which runs all models, regardless of whether an output file for the model exists, <code>"never"</code> , which does not run any model that has an existing output file, and <code>"modified"</code> , which only runs a model if the modified date for the input file is more recent than the output file modified date.

print	a character string indicating which results to show, i.e. "all" (default) for all results "input" for input command sections, and "result" for result sections.
input	a character vector specifying Mplus input command sections included in the output (see 'Details' in the mplus.print function).
result	a character vector specifying Mplus result sections included in the output (see 'Details' in the mplus.print function).
exclude	a character vector specifying Mplus input command or result sections excluded from the output (see 'Details' in the mplus.print function).
variable	logical: if TRUE, names of the variables in the data set (NAMES ARE) specified in the VARIABLE: command section are shown. By default, names of the variables in the data set are excluded from the output unless all variables are used in the analysis (i.e., no USEVARIABLES command specified in the Mplus input file).
not.input	logical: if TRUE (default), character vector indicating the input commands not requested are shown on the console.
not.result	logical: if TRUE (default), character vector indicating the result sections not requested are shown on the console.
write	a character string naming a file for writing the output into a text file with file extension ".txt" (e.g., "Output.txt").
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console by using the function mplus.print .

Details

Mplus Input Sections The function is used to update following Mplus input sections:

- TITLE
- DATA
- DATA IMPUTATION
- DATA WIDETOLONG
- DATA LONGTOWIDE
- DATA TWOPARTE
- DATA MISSING
- DATA SURVIVAL
- DATA COHORT
- VARIABLE
- DEFINE
- ANALYSIS
- MODEL
- MODEL INDIRECT
- MODEL CONSTRAINT
- MODEL TEST

- MODEL PRIORS
- MODEL MONTECARLO
- MODEL POPULATION
- MODEL COVERAGE
- MODEL MISSING
- OUTPUT
- SAVEDATA
- PLOT

The ...; Specification The ...; Specification is used to update specific options in the VARIABLE and ANALYSIS section, while keeping all other options in the `misty.object` of type `mplus` specified in the argument `x`. The ...; specification is only available for the VARIABLE and ANALYSIS section. Note that ...; including the semicolon ; needs to be specified, i.e., ... without the semicolon ; will result in an error message.

The --; Specification The --; specification is used to remove entire sections (e.g., OUTPUT: --;) or options within the VARIABLE: and ANALYSIS: section (e.g., ANALYSIS: ESTIMATOR IS --;) from the Mplus input. Note that --; including the semicolon ; needs to be specified, i.e., -- without the semicolon ; will result in an error message.

Comments in the Mplus Input Comments in the Mplus Input can cause problems when following keywords in uppercase, lower case, or mixed upper and lower case letters are involved in the comments of the VARIABLE and ANALYSIS section:

- **VARIABLE section:** "NAMES", "USEOBSERVATIONS", "USEVARIABLES", "MISSING", "CENSORED", "CATEGORICAL", "NOMINAL", "COUNT", "DSURVIVAL", "GROUPING", "IDVARIABLE", "FREQWEIGHT", "TSCORES", "AUXILIARY", "CONSTRAINT", "PATTERN", "STRATIFICATION", "CLUSTER", "WEIGHT", "WTSCALE", "BWEIGHT", "B2WEIGHT", "B3WEIGHT", "BWTSCALE", "REPWEIGHTS", "SUBPOPULATION", "FINITE", "CLASSES", "KNOWNCLASS", "TRAINING", "WITHIN", "BETWEEN", "SURVIVAL", "TIMECENSORED", "LAGGED", or "TINTERVAL".
- **ANALYSIS section:** "TYPE", "ESTIMATOR", "MODEL", "ALIGNMENT", "DISTRIBUTION", "PARAMETERIZATION", "LINK", "ROTATION", "ROWSTANDARDIZATION", "PARALLEL", "REPSE", "BASEHAZARD", "CHOLESKY", "ALGORITHM", "INTEGRATION", "MCSEED", "ADAPTIVE", "INFORMATION", "BOOTSTRAP", "LRTBOOTSTRAP", "STARTS", "STITERATIONS", "STCONVERGENCE", "STSCALE", "STSEED", "OPTSEED", "K-1STARTS", "LRTSTARTS", "RSTARTS", "ASTARTS", "H1STARTS", "DIFFTEST", "MULTIPLIER", "COVERAGE", "ADDFREQUENCY", "ITERATIONS", "SDITERATIONS", "H1ITERATIONS", "MITERATIONS", "MCITERATIONS", "MUITERATIONS", "RITERATIONS", "AITERATIONS", "CONVERGENCE", "H1CONVERGENCE", "LOGCRITERION", "RLOGCRITERION", "MCONVERGENCE", "MCCONVERGENCE", "MUCONVERGENCE", "RCONVERGENCE", "ACONVERGENCE", "MIXC", "MIXU", "LOGHIGH", "LOGLOW", "UCELLSIZE", "VARIANCE", "SIMPLICITY", "TOLERANCE", "METRIC", "MATRIX", "POINT", "CHAINS", "BSEED", "STVALUES", "PREDICTOR", "ALGORITHM", "BCONVERGENCE", "BITERATIONS", "FBITERATIONS", "THIN", "MDITERATIONS", "KOLMOGOROV", "PRIOR", "INTERACTIVE", or "PROCESSORS".

Note that comments are removed from the input text by default, i.e., `comment = FALSE`.

Value

Returns an object of class `misty.object`, which is a list with following entries:

`call` function call

type	type of analysis
x	misty.object object of type mplus
args	specification of function arguments
input	list with input command sections
write	updated write command sections
result	list with input command sections (input) and result sections (input)

Author(s)

Takuya Yanagida

References

Muthen, L. K., & Muthen, B. O. (1998-2017). *Mplus User's Guide* (8th ed.). Muthen & Muthen.

See Also

[read.mplus](#), [write.mplus](#), [mplus](#), [mplus.print](#), [mplus.plot](#), [mplus.bayes](#), [mplus.run](#), [mplus.lca](#)

Examples

```
## Not run:

#-----
# Example 1: Update VARIABLE and MODEL section

# Write Mplus Data File
write.mplus(ex3_1, file = "ex3_1.dat")

# Specify Mplus input
input <- '
DATA:      FILE IS ex3_1.dat;
VARIABLE:  NAMES ARE y1 x1 x3;
MODEL:     y1 ON x1 x3;
OUTPUT:    SAMPSTAT;
'

# Run Mplus input
mod0 <- mplus(input, file = "ex3_1.inp")

# Update VARIABLE and MODEL section
update1 <- '
VARIABLE:  ...;
           USEVARIABLES ARE y1 x1;
MODEL:     y1 ON x1;
'

# Run updated Mplus input
mod1 <- mplus.update(mod0, update1, file = "ex3_1_update1.inp")
```

```

#-----
# Example 2: Update ANALYSIS section

# Update ANALYSIS section
update2 <- '
ANALYSIS: ESTIMATOR IS MLR;
'

# Run updated Mplus input
mod2 <- mplus.update(mod1, update2, file = "ex3_1_update2.inp")

#-----
# Example 3: Remove OUTPUT section

# Remove OUTPUT section
update3 <- '
OUTPUT: ---;
'

# Run updated Mplus input
mod3 <- mplus.update(mod2, update3, file = "ex3_1_update3.inp")

## End(Not run)

```

multilevel.cfa

Multilevel Confirmatory Factor Analysis

Description

This function conducts multilevel confirmatory factor analysis to investigate four types of constructs, i.e., within-cluster constructs, shared cluster-level constructs, configural cluster constructs, and simultaneous shared and configural cluster constructs by calling the `cfa` function in the R package **lavaan**. By default, the function specifies and estimates a configural cluster construct and provides a table with univariate sample statistics, model fit information, and parameter estimates. Additionally, variance-covariance coverage of the data, modification indices, and residual correlation matrix can be requested by specifying the argument `print`.

Usage

```

multilevel.cfa(data, ..., cluster, model = NULL, rescov = NULL,
               model.w = NULL, model.b = NULL, rescov.w = NULL, rescov.b = NULL,
               const = c("within", "shared", "config", "shareconf"),
               fix.resid = NULL, ident = c("marker", "var", "effect"),
               ls.fit = FALSE, estimator = c("ML", "MLR"),
               optim.method = c("nlminb", "em"), missing = c("listwise", "fiml"),
               print = c("all", "summary", "coverage", "descript", "fit", "est",
                         "modind", "resid"),
               mod.minval = 6.63, resid.minval = 0.1, digits = 3, p.digits = 3,
               as.na = NULL, write = NULL, append = TRUE, check = TRUE, output = TRUE)

```

Arguments

<code>data</code>	a data frame. If <code>model</code> , <code>model.w</code> , and <code>model.b</code> are NULL, multilevel confirmatory factor analysis based on a measurement model with one factor labeled <code>wf</code> at the Within level and one factor labeled <code>bf</code> at the Between level comprising all variables in the data frame is conducted. Note that the cluster variable specified in <code>cluster</code> is excluded from data when specifying the argument <code>cluster</code> using the variable name of the cluster variable. If <code>model</code> or <code>model.w</code> and <code>model.b</code> is specified, the data frame needs to contain all variables used in the <code>model</code> argument(s).
<code>...</code>	an expression indicating the variable names in data. Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>cluster</code>	either a character string indicating the variable name of the cluster variable in data or data, or a vector representing the nested grouping structure (i.e., group or cluster variable).
<code>model</code>	a character vector for specifying the same factor structure with one factor at the Within and Between Level, or a list of character vectors for specifying the same measurement model with more than one factor at the Within and Between Level, e.g., <code>model = c("x1", "x2", "x3", "x4")</code> for specifying a measurement model with one factor labeled <code>wf</code> at the Within level and a measurement model with one factor labeled <code>bf</code> at the Between level each comprising four indicators, or <code>model = list(factor1 = c("x1", "x2", "x3", "x4"), factor2 = c("x5", "x6", "x7", "x8"))</code> for specifying a measurement model with two latent factors labeled <code>wfactor1</code> and <code>wfactor2</code> at the Within level and a measurement model with two latent factors labeled <code>bfactor1</code> and <code>bfactor2</code> at the Between level each comprising four indicators. Note that the name of each list element is used to label factors, where prefixes <code>w</code> and <code>b</code> are added the labels to distinguish factor labels at the Within and Between level, i.e., all list elements need to be named, otherwise factors are labeled with <code>"wf1"</code> , <code>"wf2"</code> , <code>"wf3"</code> for labels at the Within level and <code>"bf1"</code> , <code>"bf2"</code> , <code>"bf3"</code> for labels at the Between level and so on.
<code>rescov</code>	a character vector or a list of character vectors for specifying residual covariances at the Within level, e.g. <code>rescov = c("x1", "x2")</code> for specifying a residual covariance between indicators <code>x1</code> and <code>x2</code> at the Within level or <code>rescov = list(c("x1", "x2"), c("x3", "x4"))</code> for specifying residual covariances between indicators <code>x1</code> and <code>x2</code> , and indicators <code>x3</code> and <code>x4</code> at the Within level. Note that residual covariances at the Between level can only be specified by using the arguments <code>model.w</code> , <code>model.b</code> , and <code>model.b</code> .
<code>model.w</code>	a character vector specifying a measurement model with one factor at the Within level, or a list of character vectors for specifying a measurement model with more than one factor at the Within level.
<code>model.b</code>	a character vector specifying a measurement model with one factor at the Between level, or a list of character vectors for specifying a measurement model with more than one factor at the Between level.
<code>rescov.w</code>	a character vector or a list of character vectors for specifying residual covariances at the Within level.

<code>rescov.b</code>	a character vector or a list of character vectors for specifying residual covariances at the Between level.
<code>const</code>	a character string indicating the type of construct(s), i.e., "within" for within-cluster constructs, "shared" for shared cluster-level constructs, "config" (default) for configural cluster constructs, and "shareconf" for simultaneous shared and configural cluster constructs.
<code>fix.resid</code>	a character vector for specifying residual variances to be fixed at 0 at the Between level, e.g., <code>fix.resid = c("x1", "x3")</code> to fix residual variances of indicators x1 and x2 at the Between level at 0. Note that it is also possible to specify <code>fix.resid = "all"</code> which fixes all residual variances at the Between level at 0 in line with the strong factorial measurement invariance assumption across cluster.
<code>ident</code>	a character string indicating the method used for identifying and scaling latent variables, i.e., "marker" for the marker variable method fixing the first factor loading of each latent variable to 1, "var" for the fixed variance method fixing the variance of each latent variable to 1, or "effect" for the effects-coding method using equality constraints so that the average of the factor loading for each latent variable equals 1.
<code>ls.fit</code>	logical: if TRUE (default) level-specific fit indices are computed when specifying a model using the arguments <code>model.w</code> and <code>model.b</code> given the model does not contain any cross-level equality constraints.
<code>estimator</code>	a character string indicating the estimator to be used: "ML" for maximum likelihood with conventional standard errors and "MLR" (default) for maximum likelihood with Huber-White robust standard errors and a scaled test statistic that is asymptotically equal to the Yuan-Bentler test statistic. Note that by default, full information maximum likelihood (FIML) method is used to deal with missing data when using "ML" (<code>missing = "fiml"</code>), whereas incomplete cases are removed listwise (i.e., <code>missing = "listwise"</code>) when using "MLR".
<code>optim.method</code>	a character string indicating the optimizer, i.e., "nlopt" (default) for the unconstrained and bounds-constrained quasi-Newton method optimizer and "em" for the Expectation Maximization (EM) algorithm.
<code>missing</code>	a character string indicating how to deal with missing data, i.e., "listwise" (default) for listwise deletion or "fiml" for full information maximum likelihood (FIML) method. Note that FIML method is only available when <code>estimator = "ML"</code> , that it takes longer to estimate the model using FIML, and that FIML is prone to convergence issues which might be resolved by switching to listwise deletion.
<code>print</code>	a character string or character vector indicating which results to show on the console, i.e. "all" for all results, "summary" for a summary of the specification of the estimation method and missing data handling in lavaan, "coverage" for the variance-covariance coverage of the data, "descript" for descriptive statistics, "fit" for model fit, "est" for parameter estimates, and "modind" for modification indices. By default, a summary of the specification, descriptive statistics, model fit, and parameter estimates are printed.
<code>mod.minval</code>	numeric value to filter modification indices and only show modifications with a modification index value equal or higher than this minimum value. By default,

	modification indices equal or higher 6.63 are printed. Note that a modification index value of 6.63 is equivalent to a significance level of $\alpha = .01$.
resid.minval	numeric value indicating the minimum absolute residual correlation coefficients and standardized means to highlight in boldface. By default, absolute residual correlation coefficients and standardized means equal or higher 0.1 are highlighted. Note that highlighting can be disabled by setting the minimum value to 1.
digits	an integer value indicating the number of decimal places to be used for displaying results. Note that loglikelihood, information criteria and chi-square test statistic is printed with digits minus 1 decimal places.
p.digits	an integer value indicating the number of decimal places to be used for displaying the <i>p</i> -value.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that as.na() function is only applied to data but not to cluster.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification, convergence and model identification is checked.
output	logical: if TRUE (default), output is shown.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	data frame used for the current analysis
args	specification of function arguments
model	specified model
model.fit	fitted lavaan object (<code>mod.fit</code>)
check	results of the convergence and model identification check
result	list with result tables, i.e., summary for the summary of the specification of the estimation method and missing data handling in lavaan, coverage for the variance-covariance coverage of the data, <code>descript</code> for descriptive statistics, <code>fit</code> for model fit, <code>est</code> for parameter estimates, and <code>modind</code> for modification indices.

Note

The function uses the functions `cfa`, `lavInspect`, `lavTech`, `modindices`, `parameterEstimates`, and `standardizedsolution` provided in the R package **lavaan** by Yves Rosseel (2012).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Rosseel, Y. (2012). lavaan: An R Package for Structural Equation Modeling. *Journal of Statistical Software*, 48, 1-36. <https://doi.org/10.18637/jss.v048.i02>

See Also

[item.cfa](#), [multilevel.fit](#), [multilevel.invar](#), [multilevel.omega](#), [multilevel.cor](#), [multilevel.descript](#)

Examples

```
## Not run:

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

#-----
# Model specification using 'data' for a one-factor model
# with the same factor structure with one factor at the Within and Between Level

#.....
# Cluster variable specification

# Example 1a: Specification using the argument '...'
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster")

# Example 1b: Alternative specification with cluster variable 'cluster' in 'data'
multilevel.cfa(Demo.twolevel[, c("y1", "y2", "y3", "y4", "cluster")], cluster = "cluster")

# Example 1c: Alternative specification with cluster variable 'cluster' not in 'data'
multilevel.cfa(Demo.twolevel[, c("y1", "y2", "y3", "y4")], cluster = Demo.twolevel$cluster)

#.....
# Type of construct

# Example 2a: Within-cluster constructs
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", const = "within")

# Example 2b: Shared cluster-level construct
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", const = "shared")

# Example 2c: Configural cluster construct (default)
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", const = "config")

# Example 2d: Simultaneous shared and configural cluster construct
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", const = "shareconf")

#.....
# Residual covariances at the Within level
```

```

# Example 3a: Residual covariance between 'y1' and 'y3'
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", rescov = c("y1", "y3"))

# Example 3b: Residual covariance between 'y1' and 'y3', and 'y2' and 'y4'
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster",
               rescov = list(c("y1", "y3"), c("y2", "y4")))

#.....
# Residual variances at the Between level fixed at 0

# Example 4a: All residual variances fixed at 0
# i.e., strong factorial invariance across clusters
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", fix.resid = "all")

# Example 4b: Residual variances of 'y1', 'y2', and 'y4' fixed at 0
# i.e., partial strong factorial invariance across clusters
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", fix.resid = c("y1", "y2", "y4"))

#.....
# Print all results

# Example 5: Set minimum value for modification indices to 1
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", print = "all",
               mod.minval = 1)

#.....
# Example 6: lavaan model and summary of the estimated model

mod <- multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", output = FALSE)

# lavaan model syntax
cat(mod$model)

# Fitted lavaan object
lavaan::summary(mod$model.fit, standardized = TRUE, fit.measures = TRUE)

#.....
# Write results

# Example 7a: Assign results into an object and write results into an Excel file
mod <- multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", print = "all",
                     write = "Multilevel_CFA.txt", output = FALSE)

# Example 7b: Assign results into an object and write results into an Excel file
mod <- multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", print = "all",
                     output = FALSE)

# Write results into an Excel file
write.result(mod, "Multilevel_CFA.xlsx")

# Estimate model and write results into an Excel file
multilevel.cfa(Demo.twolevel, y1:y4, cluster = "cluster", print = "all",

```

```

write = "Multilevel_CFA.xlsx")

#-----
# Model specification using 'model' for one or multiple factor model
# with the same factor structure at the Within and Between Level

# Example 8a: One-factor model
multilevel.cfa(Demo.twolevel, cluster = "cluster", model = c("y1", "y2", "y3", "y4"))

# Example 8b: Two-factor model
multilevel.cfa(Demo.twolevel, cluster = "cluster",
               model = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")))

# Example 8c: Two-factor model with user-specified labels for the factors
multilevel.cfa(Demo.twolevel, cluster = "cluster",
               model = list(factor1 = c("y1", "y2", "y3"), factor2 = c("y4", "y5", "y6")))

#.....
# Type of construct

# Example 9a: Within-cluster constructs
multilevel.cfa(Demo.twolevel, cluster = "cluster", const = "within",
               model = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")))

# Example 9b: Shared cluster-level construct
multilevel.cfa(Demo.twolevel, cluster = "cluster", const = "shared",
               model = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")))

# Example 9c: Configural cluster construct (default)
multilevel.cfa(Demo.twolevel, cluster = "cluster", const = "config",
               model = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")))

# Example 9d: Simultaneous shared and configural cluster construct
multilevel.cfa(Demo.twolevel, cluster = "cluster", const = "shareconf",
               model = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")))

#.....
# Residual covariances at the Within level

# Example 10a: Residual covariance between 'y1' and 'y4' at the Within level
multilevel.cfa(Demo.twolevel, cluster = "cluster",
               model = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")),
               rescov = c("y1", "y4"))

# Example 10b: Fix all residual variances at 0
# i.e., strong factorial invariance across clusters
multilevel.cfa(Demo.twolevel, cluster = "cluster",
               model = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")),
               fix.resid = "all")

#-----
# Model specification using 'model.w' and 'model.b' for one or multiple factor model
# with different factor structure at the Within and Between Level

```

```
# Example 11a: Two-factor model at the Within level and one-factor model at the Between level
multilevel.cfa(Demo.twolevel, cluster = "cluster",
               model.w = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")),
               model.b = c("y1", "y2", "y3", "y4", "y5", "y6"))

# Example 11b: Residual covariance between 'y1' and 'y4' at the Within level
# Residual covariance between 'y5' and 'y6' at the Between level
multilevel.cfa(Demo.twolevel, cluster = "cluster",
               model.w = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")),
               model.b = c("y1", "y2", "y3", "y4", "y5", "y6"),
               rescov.w = c("y1", "y4"),
               rescov.b = c("y5", "y6"))

## End(Not run)
```

multilevel.cor

Within-Group and Between-Group Correlation Matrix

Description

This function computes the within-group and between-group correlation matrix by calling the `sem` function in the R package **lavaan** and provides standard errors, z test statistics, and significance values (*p*-values) for testing the hypothesis $H_0: \rho = 0$ for all pairs of variables within and between groups. By default, the function computes the within-group and between-group correlation matrix without standard errors, z test statistics, and significance value.

Usage

```
multilevel.cor(data, ..., cluster, within = NULL, between = NULL,
               estimator = c("ML", "MLR"), optim.method = c("nlminb", "em"),
               optim.switch = TRUE, missing = c("listwise", "fiml"), sig = FALSE,
               alpha = 0.05, print = c("all", "cor", "se", "stat", "p"), split = FALSE,
               order = FALSE, tri = c("both", "lower", "upper"), tri.lower = TRUE,
               p.adj = c("none", "bonferroni", "holm", "hochberg", "hommel",
                          "BH", "BY", "fdr"), digits = 2, p.digits = 3,
               as.na = NULL, write = NULL, append = TRUE, check = TRUE,
               output = TRUE)
```

Arguments

<code>data</code>	a data frame.
<code>...</code>	an expression indicating the variable names in data, e.g., <code>multilevel.cor(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>cluster</code>	either a character string indicating the variable name of the cluster variable in data, or a vector representing the nested grouping structure (i.e., group or cluster variable).

within	a character vector representing variables that are measured on the within level and modeled only on the within level. Variables not mentioned in within or between are measured on the within level and will be modeled on both the within and between level.
between	a character vector representing variables that are measured on the between level and modeled only on the between level. Variables not mentioned in within or between are measured on the within level and will be modeled on both the within and between level.
estimator	a character string indicating the estimator to be used, i.e., "ML" for maximum likelihood with conventional standard errors and "MLR" for maximum likelihood with Huber-White robust standard errors. The default setting depends on the argument sig, i.e., "ML" is used when specifying sig = FALSE (default) and "MLR" is used when specifying sig = TRUE.
optim.method	a character string indicating the optimizer, i.e., "nlminb" (default) for the unconstrained and bounds-constrained quasi-Newton method optimizer and "em" for the Expectation Maximization (EM) algorithm.
optim.switch	logical: if TRUE (default), model estimation switches to Expectation Maximization (EM) algorithm ("em") if the quasi-Newton optimization ("nlminb" (default)) does not converge.
missing	a character string indicating how to deal with missing data, i.e., "listwise" for listwise deletion or "fiml" (default) for full information maximum likelihood (FIML) method. Note that it takes longer to estimate models while using FIML and using FIML is prone to issues with model convergence, these issues might be resolved by switching to listwise deletion.
sig	logical: if TRUE, statistically significant correlation coefficients are shown in boldface on the console. Note that standard errors, z test statistics, and significance values not provided in the return object when sig = FALSE (default).
alpha	a numeric value between 0 and 1 indicating the significance level at which correlation coefficients are printed boldface when sig = TRUE.
print	a character string or character vector indicating which results to show on the console, i.e. "all" for all results, "cor" for correlation coefficients, "se" for standard errors, "stat" for z test statistics, and "p" for <i>p</i> -values.
split	logical: if TRUE, output table is split in within-group and between-group correlation matrix.
order	logical: if TRUE, variables in the output table are ordered, so that variables specified in the argument between are shown first.
tri	a character string indicating which triangular of the matrix to show on the console when split = TRUE, i.e., both for upper and upper for the upper triangular.
tri.lower	logical: if TRUE (default) and split = FALSE (default), within-group correlations are shown in the lower triangular and between-group correlation are shown in the upper triangular.
p.adj	a character string indicating an adjustment method for multiple testing based on p.adjust , i.e., none (default), bonferroni, holm, hochberg, hommel, BH, BY, or fdr.

<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying correlation coefficients.
<code>p.digits</code>	an integer value indicating the number of decimal places to be used for displaying p -values.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to data but not to <code>cluster</code> .
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

Within-Group and Between-Group Variables The specification of the within-group and between-group variables is in line with the syntax in Mplus. That is, the `within` argument is used to identify variables in the data frame specified in `data` that are measured at the individual level and modeled only at the within level. They are specified to have no variance in the between part of the model. The `between` argument is used to identify the variables in the data frame specified in `data` that are measured at the cluster level and modeled only at the between level. Variables not mentioned in the arguments `within` or `between` are measured at the individual level and will be modeled at both the within and between level.

Estimation Method and Missing Data Handling The default setting for the argument `estimator` is depending on the setting of the argument `sig`. If `sig = FALSE` (default), maximum likelihood estimation (`estimator = "ML"`) is used, while maximum likelihood with Huber-White robust standard errors (`estimator = "MLR"`) that are robust against non-normality is used when `sig = TRUE`. In the presence of missing data, full information maximum likelihood (FIML) method (`missing = "fiml"`) is used by default. Note that FIML method cannot deal with within-group variables that have no variance within some clusters. In this cases, the function will switch to listwise deletion. Using FIML method might result in issues with model convergence, which will be resolved by switching to listwise deletion (`missing = "listwise"`).

Optimizer The lavaan package uses a quasi-Newton optimization method ("`nlminb`") by default. If the optimizer does not converge, model estimation switches to the Expectation Maximization (EM) algorithm ("`em`") if the argument `optim.switch` is specified as TRUE (default).

Statistical Significance Statistically significant correlation coefficients can be shown in boldface on the console by specifying `sig = TRUE`. However, this option is not supported when using R Markdown, i.e., the argument `sig` will switch to FALSE.

Adjustment Method for Multiple Testing Adjustment method for multiple testing when specifying the argument `p.adj` is applied to the within-group and between-group correlation matrix separately.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame specified in <code>data</code> including the group variable specified in <code>cluster</code>
<code>args</code>	specification of function arguments
<code>model.fit</code>	fitted lavaan object (<code>mod.fit</code>)
<code>result</code>	list with result tables, i.e., summary for the specification of the estimation method and missing data handling in lavaan, <code>wb.cor</code> for the within- and between-group correlations, <code>wb.se</code> for the standard error of the within- and between-group correlations, <code>wb.stat</code> for the test statistic of within- and between-group correlations, <code>wb.p</code> for the significance value of the within- and between-group correlations, <code>with.cor</code> for the within-group correlations, <code>with.se</code> for the standard error of the within-group correlations, <code>with.stat</code> for the test statistic of within-group correlations, <code>with.p</code> for the significance value of the within-group correlations, <code>betw.cor</code> for the between-group correlations, <code>betw.se</code> for the standard error of the between-group correlations, <code>betw.stat</code> for the test statistic of between-group correlations, <code>betw.p</code> for the significance value of the between-group correlations

Note

The function uses the functions `sem`, `lavInspect`, `lavMatrixRepresentation`, `lavTech`, `parameterEstimates`, and `standardizedsolution` provided in the R package **lavaan** by Yves Rosseel (2012).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Hox, J., Moerbeek, M., & van de Schoot, R. (2018). *Multilevel analysis: Techniques and applications* (3rd. ed.). Routledge.
- Snijders, T. A. B., & Bosker, R. J. (2012). *Multilevel analysis: An introduction to basic and advanced multilevel modeling* (2nd ed.). Sage Publishers.

See Also

[write.result](#), [multilevel.descript](#), [multilevel.icc](#), [cluster.scores](#)

Examples

```
## Not run:

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")
```

```

#-----
# Cluster variable specification

# Example 1: Specification using the argument '...'
multilevel.cor(Demo.twolevel, y1, y2, y3, cluster = "cluster")

# Alternative specification with cluster variable 'cluster' in 'data'
multilevel.cor(Demo.twolevel[, c("y1", "y2", "y3", "cluster")], cluster = "cluster")

# Alternative specification with cluster variable 'cluster' not in 'data'
multilevel.cor(Demo.twolevel[, c("y1", "y2", "y3")], cluster = Demo.twolevel$cluster)

#-----
# Example 2: All variables modeled at both the within and between level
# Highlight statistically significant result at alpha = 0.05
multilevel.cor(Demo.twolevel, y1, y2, y3, sig = TRUE, cluster = "cluster")

# Example 3: Split output table in within-group and between-group correlation matrix.
multilevel.cor(Demo.twolevel, y1, y2, y3, cluster = "cluster", split = TRUE)

# Example 4: Print correlation coefficients, standard errors, z test statistics,
# and p-values
multilevel.cor(Demo.twolevel, y1, y2, y3, cluster = "cluster", sig = TRUE, print = "all")

# Example 5: Print correlation coefficients and p-values
# significance values with Bonferroni correction
multilevel.cor(Demo.twolevel, y1, y2, y3, cluster = "cluster", sig = TRUE,
               print = c("cor", "p"), p.adj = "bonferroni")

#-----
# Example 6: Variables "y1", "y2", and "y2" modeled at both the within and between level
# Variables "w1" and "w2" modeled at the cluster level
multilevel.cor(Demo.twolevel, y1, y2, y3, w1, w2, cluster = "cluster",
               between = c("w1", "w2"))

# Example 7: Show variables specified in the argument 'between' first
multilevel.cor(Demo.twolevel, y1, y2, y3, w1, w2, cluster = "cluster",
               between = c("w1", "w2"), order = TRUE)

#-----
# Example 8: Variables "y1", "y2", and "y2" modeled only at the within level
# Variables "w1" and "w2" modeled at the cluster level
multilevel.cor(Demo.twolevel, y1, y2, y3, w1, w2, cluster = "cluster",
               within = c("y1", "y2", "y3"), between = c("w1", "w2"))

#-----
# Example 9: lavaan model and summary of the multilevel model used to compute the
# within-group and between-group correlation matrix

mod <- multilevel.cor(Demo.twolevel, y1, y2, y3, cluster = "cluster", output = FALSE)

# lavaan model syntax
mod$model

```

```

# Fitted lavaan object
lavaan::summary(mod$model.fit, standardized = TRUE)

#-----
# Write Results

# Example 10a: Write Results into a text file
multilevel.cor(Demo.twolevel, y1, y2, y3, cluster = "cluster",
               write = "Multilevel_Correlation.txt")

# Example 10b: Write Results into a Excel file
multilevel.cor(Demo.twolevel, y1, y2, y3, cluster = "cluster",
               write = "Multilevel_Correlation.xlsx")

## End(Not run)

```

multilevel.descript *Multilevel Descriptive Statistics for Two-Level and Three-Level Data*

Description

This function computes descriptive statistics for two-level and three-level multilevel data, e.g. average cluster size, variance components, intraclass correlation coefficient, design effect, and effective sample size.

Usage

```

multilevel.descript(data, ..., cluster, type = c("1a", "1b"),
                    method = c("aov", "lme4", "nlme"),
                    print = c("all", "var", "sd"), REML = TRUE, digits = 2,
                    icc.digits = 3, as.na = NULL, write = NULL, append = TRUE,
                    check = TRUE, output = TRUE)

```

Arguments

data	a numeric vector or data frame.
...	an expression indicating the variable names in data. Note that the operators +, -, ~, :, ::, and ! can also be used to select variables, see 'Details' in the df.subset function.
cluster	a character string indicating the name of the cluster variable in data for two-level data, a character vector indicating the names of the cluster variables in data for three-level data, or a vector or data frame representing the nested grouping structure (i.e., group or cluster variables). Alternatively, a character string or character vector indicating the variable name(s) of the cluster variable(s) in data. Note that the cluster variable at Level 3 come first in a three-level model, i.e., cluster = c("level3", "level2").

type	a character string indicating the type of intraclass correlation coefficient, i.e., type = "1a" (default) for ICC(1) representing the proportion of variance at Level 2 and Level 3, type = "1b" representing an estimate of the expected correlation between two randomly chosen elements in the same group when specifying a three-level model (i.e., two cluster variables). See 'Details' in the multilevel.icc function for the formula used in this function.
method	a character string indicating the method used to estimate intraclass correlation coefficients, i.e., "aov" ICC estimated using the aov function, "lme4" (default) ICC estimated using the lmer function in the lme4 package, "nlme" ICC estimated using the lme function in the nlme package.
print	a character string or character vector indicating which results to show on the console, i.e. "all" for variances and standard deviations, "var" (default) for variances, or "sd" for standard deviations within and between clusters.
REML	logical: if TRUE (default), restricted maximum likelihood is used to estimate the null model when using the lmer() function in the lme4 package or the lme() function in the nlme package.
digits	an integer value indicating the number of decimal places to be used.
icc.digits	an integer indicating the number of decimal places to be used for displaying intraclass correlation coefficients.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that as.na() function is only applied to data but not to cluster.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

Two-Level Model In a two-level model, the intraclass correlation coefficients, design effect, and the effective sample size are computed based on the random intercept-only model:

$$Y_{ij} = \gamma_{00} + u_{0j} + r_{ij}$$

where the variance in Y is decomposed into two independent components: $\sigma_{u_0}^2$, which represents the variance at Level 2, and σ_r^2 , which represents the variance at Level 1 (Hox et al., 2018). For the computation of the intraclass correlation coefficients, see 'Details' in the [multilevel.icc](#) function. The design effect represents the effect of cluster sampling on the variance of parameter estimation and is defined by the equation

$$def = \left(\frac{SE_{Cluster}}{SE_{Simple}} \right)^2 = 1 + \rho(J - 1)$$

where $SE_{Cluster}$ is the standard error under cluster sampling, SE_{Simple} is the standard error under simple random sampling, ρ is the intraclass correlation coefficient, ICC(1), and J is the average cluster size. The effective sample size is defined by the equation:

$$N_{effective} = \frac{N_{total}}{deff}$$

The effective sample size $N_{effective}$ represents the equivalent total sample size that we should use in estimating the standard error (Snijders & Bosker, 2012).

Three-Level Model In a three-level model, the intraclass correlation coefficients, design effect, and the effective sample size are computed based on the random intercept-only model:

$$Y_{ijk} = \gamma_{000} + v_{0k} + u_{0jk} + r_{ijk}$$

where the variance in Y is decomposed into three independent components: $\sigma_{v_0}^2$, which represents the variance at Level 3, $\sigma_{u_0}^2$, which represents the variance at Level 2, and σ_r^2 , which represents the variance at Level 1 (Hox et al., 2018). For the computation of the intraclass correlation coefficients, see 'Details' in the `multilevel.icc` function. The design effect represents the effect of cluster sampling on the variance of parameter estimation and is defined by the equation

$$deff = \left(\frac{SE_{Cluster}}{SE_{Simple}} \right)^2 = 1 + \rho_{L2}(J - 1) + \rho_{L3}(JK - 1)$$

where ρ_{L2} is the ICC(1) at Level 2, ρ_{L3} is the ICC(1) at Level 3, J is the average cluster size at Level 2, and K is the average cluster size at Level 3.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame specified in <code>data</code> including the cluster variable(s) specified in <code>cluster</code>
<code>args</code>	specification of function arguments
<code>model.fit</code>	fitted lavaan object (<code>mod.fit</code>)
<code>result</code>	list with result tables, i.e., <code>no.obs</code> for the number of observations, <code>no.no.miss</code> for the number of missing value, <code>no.cluster.l2</code> and <code>no.cluster.l3</code> for the number of clusters at Level 2 and/or Level 3, <code>m.cluster.size.l2</code> and <code>m.cluster.size.l3</code> for the average cluster size at Level 2 and/or Level 3, <code>sd.cluster.size.l2</code> and <code>sd.cluster.size.l3</code> for the standard deviation of the cluster size at Level 2 and/or Level 3, <code>min.cluster.size.l2</code> <code>min.cluster.size.l3</code> for the minimum cluster size at Level 2 and/or Level 3, <code>max.cluster.size.l2</code> <code>max.cluster.size.l3</code> for the maximum cluster size at Level 2 and/or Level 3, <code>mean.x</code> for the intercept of the multilevel model, <code>var.r</code> for the variance within clusters, <code>var.u</code> for the variance between Level 2 clusters, <code>var.b</code> for the variance between Level 3 clusters, <code>icc1.l2</code> and <code>icc1.l3</code> for ICC(1) at Level 2 and/or Level 3, <code>icc2.l2</code> and <code>icc2.l3</code> for ICC(2) at Level 2 and/or Level 3, <code>deff</code> for the design effect, <code>deff.sqrt</code> for the square root of the design effect, <code>n.effect</code> for the effective sample size

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Hox, J., Moerbeek, M., & van de Schoot, R. (2018). *Multilevel analysis: Techniques and applications* (3rd. ed.). Routledge.
- Snijders, T. A. B., & Bosker, R. J. (2012). *Multilevel analysis: An introduction to basic and advanced multilevel modeling* (2nd ed.). Sage Publishers.

See Also

[write.result](#), [multilevel.icc](#), [descript](#)

Examples

```
## Not run:

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

#-----
# Two-Level Data

#.....
# Cluster variable specification

# Example 1a: Specification using the argument '...'
multilevel.descript(Demo.twolevel, y1, cluster = "cluster")

# Example 1b: Alternative specification with cluster variable 'cluster' in 'data'
multilevel.descript(Demo.twolevel[, c("y1", "cluster")], cluster = "cluster")

# Example 1c: Alternative specification with cluster variable 'cluster' not in 'data'
multilevel.descript(Demo.twolevel$y1, cluster = Demo.twolevel$cluster)

#-----

# Example 2: Multilevel descriptive statistics for 'y1'
multilevel.descript(Demo.twolevel, y1, cluster = "cluster")

# Example 3: Multilevel descriptive statistics, print variance and standard deviation
multilevel.descript(Demo.twolevel, y1, cluster = "cluster", print = "all")

# Example 4: Multilevel descriptive statistics, print ICC with 5 digits
multilevel.descript(Demo.twolevel, y1, cluster = "cluster", icc.digits = 5)

# Example 5: Multilevel descriptive statistics
# use lme() function in the nlme package to estimate ICC
multilevel.descript(Demo.twolevel, y1, cluster = "cluster", method = "nlme")
```

```

# Example 6a: Multilevel descriptive statistics for 'y1', 'y2', 'y3', 'w1', and 'w2'
multilevel.descript(Demo.twolevel, y1, y2, y3, w1, w2, cluster = "cluster")

# Alternative specification without using the '...' argument
multilevel.descript(Demo.twolevel[, c("y1", "y2", "y3", "w1", "w2")],
                    cluster = Demo.twolevel$cluster)

#-----
# Three-Level Data

# Create arbitrary three-level data
Demo.threelevel <- data.frame(Demo.twolevel, cluster2 = Demo.twolevel$cluster,
                             cluster3 = rep(1:10, each = 250))

#.....
# Cluster variable specification

# Example 7a: Specification using the argument '...'
multilevel.descript(Demo.threelevel, y1, cluster = c("cluster3", "cluster2"))

# Example 7b: Alternative specification without using the argument '...'
multilevel.descript(Demo.threelevel[, c("y1", "cluster3", "cluster2")],
                    cluster = c("cluster3", "cluster2"))

# Example 7c: Alternative specification with cluster variables 'cluster' not in 'data'
multilevel.descript(Demo.threelevel$y1,
                    cluster = Demo.threelevel[, c("cluster3", "cluster2")])

#-----

# Example 8: Multilevel descriptive statistics for 'y1', 'y2', 'y3', 'w1', and 'w2'
multilevel.descript(Demo.threelevel, y1:y3, w1, w2, cluster = c("cluster3", "cluster2"))

#-----
# Write Results

# Example 9a: Write Results into a text file
multilevel.descript(Demo.twolevel, y1, y2, y3, w1, w2, cluster = "cluster",
                    write = "Multilevel_Descript.txt")

# Example 9b: Write Results into a Excel file
multilevel.descript(Demo.twolevel, y1, y2, y3, w1, w2, cluster = "cluster",
                    write = "Multilevel_Descript.xlsx")

## End(Not run)

```

Description

This function provides simultaneous and level-specific model fit information using the partially saturated model method for multilevel models estimated with the **lavaan** package. Note that level-specific fit indices cannot be computed when the fitted model contains cross-level constraints, e.g., equal factor loadings across levels in line with the metric cross-level measurement invariance assumption.

Usage

```
multilevel.fit(model, print = c("all", "summary", "fit"), digits = 3, p.digits = 3,
               write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

model	a fitted model of class "lavaan" from the lavaan package.
print	a character string or character vector indicating which results to show on the console, i.e. "all" for all results, "summary" for a summary of the specification of the estimation method and missing data handling in lavaan and "fit" for model fit.
digits	an integer value indicating the number of decimal places to be used for displaying results. Note that loglikelihood, information criteria and chi-square test statistic is printed with digits minus 1 decimal places.
p.digits	an integer value indicating the number of decimal places to be used for displaying the <i>p</i> -value.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
model	a fitted model of class "lavaan"
args	specification of function arguments
model	specified models, i.e., <code>mod.11</code> for the model at the Within level, <code>mod.11.syntax</code> for the lavaan syntax for the model at the Between level, <code>mod.12</code> for the model at the Within level, <code>mod.12.syntax</code> for the lavaan syntax for the model at the Between level, <code>mod.112</code> for the model at the Within and Between level,

	mod.l12.syntax for the lavaan syntax for the model at the Within and Between level, l1.mod.base for the baseline model at the Within level saturated at the Between level, l1.mod.hypo for the hypothesized model at the Within level saturated at the Between level, l2.mod.base for the baseline model at the Between level saturated at the Within level, l2.mod.hypo for the hypothesized model at the Between level saturated at the Within level
result	list with result tables, i.e., summary for the summary of the specification of the estimation method and missing data handling in lavaan and fit for the model fit information.

Note

The function uses the functions `cfa`, `fitmeasures`, `lavInspect`, `lavTech`, and `parTable` provided in the R package **lavaan** by Yves Rosseel (2012).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Rosseel, Y. (2012). lavaan: An R Package for Structural Equation Modeling. *Journal of Statistical Software*, 48, 1-36. <https://doi.org/10.18637/jss.v048.i02>

See Also

[multilevel.cfa](#), [multilevel.invar](#), [multilevel.omega](#), [multilevel.cor](#), [multilevel.descript](#)

Examples

```
## Not run:

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

# Model specification
model <- 'level: 1
          fw =~ y1 + y2 + y3
          fw ~ x1 + x2 + x3
          level: 2
          fb =~ y1 + y2 + y3
          fb ~ w1 + w2'

#-----

# Example 1: Model estimation with estimator = "ML"
fit1 <- lavaan::sem(model = model, data = Demo.twolevel, cluster = "cluster",
                   estimator = "ML")

# Simultaneous and level-specific multilevel model fit information
ls.fit1 <- multilevel.fit(fit1)
```

```
# Write results into a text file
multilevel.fit(fit1, write = "LS-Fit1.txt")

# Write results into an Excel file
write.result(ls.fit1, "LS-Fit1.xlsx")

# Example 2: Model estimation with estimator = "MLR"
fit2 <- lavaan::sem(model = model, data = Demo.twolevel, cluster = "cluster",
                    estimator = "MLR")

# Simultaneous and level-specific multilevel model fit information
# Write results into an Excel file
multilevel.fit(fit2, write = "LS-Fit2.xlsx")

## End(Not run)
```

multilevel.icc	<i>Intraclass Correlation Coefficient, ICC(1) and ICC(2)</i>
----------------	--

Description

This function computes the intraclass correlation coefficient ICC(1), i.e., proportion of the total variance explained by the grouping structure, and ICC(2), i.e., reliability of aggregated variables in a two-level and three-level model.

Usage

```
multilevel.icc(data, ..., cluster, type = c("1a", "1b", "2"),
               method = c("aov", "lme4", "nlme"), REML = TRUE,
               as.na = NULL, check = TRUE)
```

Arguments

data	a numeric vector or data frame.
...	an expression indicating the variable names in data. Note that the operators +, -, ~, :, ::, and ! can also be used to select variables, see 'Details' in the df.subset function.
cluster	a character string indicating the name of the cluster variable in data for two-level data, a character vector indicating the names of the cluster variables in data for three-level data, or a vector or data frame representing the nested grouping structure (i.e., group or cluster variables). Alternatively, a character string or character vector indicating the variable name(s) of the cluster variable(s) in data. Note that the cluster variable at Level 3 come first in a three-level model, i.e., cluster = c("level3", "level2").

type	a character string indicating the type of intraclass correlation coefficient, i.e., type = "1a" (default) for ICC(1) and type = "2" for ICC(2) when specifying a two-level model (i.e., one cluster variable), and type = "1a" (default) for ICC(1) representing the proportion of variance at Level 2 and Level 3, type = "1b" representing an estimate of the expected correlation between two randomly chosen elements in the same group, and type = "2" for ICC(2) when specifying a three-level model (i.e., two cluster variables). See 'Details' for the formula used in this function.
method	a character string indicating the method used to estimate intraclass correlation coefficients, i.e., method = "aov" ICC estimated using the aov function, method = "lme4" (default) ICC estimated using the lmer function in the lme4 package, method = "nlme" ICC estimated using the lme function in the nlme package. Note that if the lme4 or nlme package is needed when estimating ICCs in a three-level model.
REML	logical: if TRUE (default), restricted maximum likelihood is used to estimate the null model when using the lmer function in the lme4 package or the lme function in the nlme package.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that as.na() function is only applied to x but not to cluster.
check	logical: if TRUE (default), argument specification is checked.

Details

Two-Level Model In a two-level model, the intraclass correlation coefficients are computed in the random intercept-only model:

$$Y_{ij} = \gamma_{00} + u_{0j} + r_{ij}$$

where the variance in Y is decomposed into two independent components: $\sigma_{u_0}^2$, which represents the variance at Level 2, and σ_r^2 , which represents the variance at Level 1 (Hox et al., 2018). These two variances sum up to the total variance and are referred to as variance components. The intraclass correlation coefficient, ICC(1) ρ requested by type = "1a" represents the proportion of the total variance explained by the grouping structure and is defined by the equation

$$\rho = \frac{\sigma_{u_0}^2}{\sigma_{u_0}^2 + \sigma_r^2}$$

The intraclass correlation coefficient, ICC(2) λ_j requested by type = "2" represents the reliability of aggregated variables and is defined by the equation

$$\lambda_j = \frac{\sigma_{u_0}^2}{\sigma_{u_0}^2 + \frac{\sigma_r^2}{n_j}} = \frac{n_j \rho}{1 + (n_j - 1)\rho}$$

where n_j is the average group size (Snijders & Bosker, 2012).

Three-Level Model In a three-level model, the intraclass correlation coefficients are computed in the random intercept-only model:

$$Y_{ijk} = \gamma_{000} + v_{0k} + u_{0jk} + r_{ijk}$$

where the variance in Y is decomposed into three independent components: $\sigma_{v_0}^2$, which represents the variance at Level 3, $\sigma_{u_0}^2$, which represents the variance at Level 2, and σ_r^2 , which represents the variance at Level 1 (Hox et al., 2018). There are two ways to compute the intraclass correlation coefficient in a three-level model. The first method requested by type = "1a" represents the proportion of variance at Level 2 and Level 3 and should be used if we are interested in a decomposition of the variance across levels. The intraclass correlation coefficient, ICC(1) ρ_{L2} at Level 2 is defined as:

$$\rho_{L2} = \frac{\sigma_{u_0}^2}{\sigma_{v_0}^2 + \sigma_{u_0}^2 + \sigma_r^2}$$

The ICC(1) ρ_{L3} at Level 3 is defined as:

$$\rho_{L3} = \frac{\sigma_{v_0}^2}{\sigma_{v_0}^2 + \sigma_{u_0}^2 + \sigma_r^2}$$

The second method requested by type = "1b" represents the expected correlation between two randomly chosen elements in the same group. The intraclass correlation coefficient, ICC(1) ρ_{L2} at Level 2 is defined as:

$$\rho_{L2} = \frac{\sigma_{v_0}^2 + \sigma_{u_0}^2}{\sigma_{v_0}^2 + \sigma_{u_0}^2 + \sigma_r^2}$$

The ICC(1) ρ_{L3} at Level 3 is defined as:

$$\rho_{L3} = \frac{\sigma_{v_0}^2}{\sigma_{v_0}^2 + \sigma_{u_0}^2 + \sigma_r^2}$$

Note that both formula are correct, but express different aspects of the data, which happen to coincide when there are only two levels (Hox et al., 2018).

The intraclass correlation coefficients, ICC(2) requested by type = "2" represent the reliability of aggregated variables at Level 2 and Level 3. The ICC(2) λ_j at Level 2 is defined as:

$$\lambda_j = \frac{\sigma_{u_0}^2}{\sigma_{u_0}^2 + \frac{\sigma_r^2}{n_j}}$$

The ICC(2) λ_k at Level 3 is defined as:

$$\lambda_k = \frac{\sigma_{v_0}^2}{\frac{\sigma_{v_0}^2 + \sigma_{u_0}^2}{n_j} + \frac{\sigma_r^2}{n_k \cdot n_j}}$$

where n_j is the average group size at Level 2 and n_j is the average group size at Level 3 (Hox et al., 2018).

Value

Returns a numeric vector or matrix with intraclass correlation coefficient(s). In a three level model, the label L2 is used for ICCs at Level 2 and L3 for ICCs at Level 3.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Hox, J., Moerbeek, M., & van de Schoot, R. (2018). *Multilevel analysis: Techniques and applications* (3rd. ed.). Routledge.

Snijders, T. A. B., & Bosker, R. J. (2012). *Multilevel analysis: An introduction to basic and advanced multilevel modeling* (2nd ed.). Sage Publishers.

See Also

[multilevel.cfa](#), [multilevel.cor](#), [multilevel.descript](#)

Examples

```
# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

#-----
# Two-Level Data

#.....
# Cluster variable specification

# Example 1a: Specification using the argument '...'
multilevel.icc(Demo.twolevel, y1, cluster = "cluster")

# Example 1b: Alternative specification with cluster variable 'cluster' in 'data'
multilevel.icc(Demo.twolevel[, c("y1", "cluster")], cluster = "cluster")

# Example 1c: Alternative specification with cluster variable 'cluster' not in 'data'
multilevel.icc(Demo.twolevel$y1, cluster = Demo.twolevel$cluster)

#.....

# Example 2: ICC(1) for 'y1'
multilevel.icc(Demo.twolevel, y1, cluster = "cluster")

# Example 3: ICC(2)
multilevel.icc(Demo.twolevel, y1, cluster = "cluster", type = "2")

# Example 4: ICC(1)
# use lme() function in the lme4 package to estimate ICC
multilevel.icc(Demo.twolevel, y1, cluster = "cluster", method = "nlme")
```

```

# Example 5: ICC(1) for 'y1', 'y2', and 'y3'
multilevel.icc(Demo.twolevel, y1, y2, y3, cluster = "cluster")

# Alternative specification without using the '...' argument
multilevel.icc(Demo.twolevel[, c("y1", "y2", "y3")], cluster = Demo.twolevel$cluster)

#-----
# Three-Level Data

# Create arbitrary three-level data
Demo.threelevel <- data.frame(Demo.twolevel, cluster2 = Demo.twolevel$cluster,
                             cluster3 = rep(1:10, each = 250))

#.....
# Cluster variable specification

# Example 6a: Specification using the argument '...'
multilevel.icc(Demo.threelevel, y1, cluster = c("cluster3", "cluster2"))

# Example 6b: Alternative specification without using the argument '...'
multilevel.icc(Demo.threelevel[, c("y1", "cluster3", "cluster2")],
               cluster = c("cluster3", "cluster2"))

# Example 6c: Alternative specification with cluster variables 'cluster' not in 'data'
multilevel.icc(Demo.threelevel$y1, cluster = Demo.threelevel[, c("cluster3", "cluster2")])

#-----

# Example 7a: ICC(1), proportion of variance at Level 2 and Level 3
multilevel.icc(Demo.threelevel, y1, cluster = c("cluster3", "cluster2"))

# Example 7b: ICC(1), expected correlation between two randomly chosen elements
# in the same group
multilevel.icc(Demo.threelevel, y1, cluster = c("cluster3", "cluster2"), type = "1b")

# Example 7c: ICC(2)
multilevel.icc(Demo.threelevel, y1, cluster = c("cluster3", "cluster2"), type = "2")

```

multilevel.indirect	<i>Confidence Interval for the Indirect Effect in a 1-1-1 Multilevel Mediation Model</i>
---------------------	--

Description

This function computes the confidence interval for the indirect effect in a 1-1-1 multilevel mediation model with random slopes based on the Monte Carlo method.

Usage

```
multilevel.indirect(a, b, se.a, se.b, cov.ab = 0, cov.rand, se.cov.rand,
```

```
nrep = 100000, alternative = c("two.sided", "less", "greater"),
seed = NULL, conf.level = 0.95, digits = 3, write = NULL,
append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>a</code>	a numeric value indicating the coefficient a , i.e., average effect of X on M on the cluster or between-group level.
<code>b</code>	a numeric value indicating the coefficient b , i.e., average effect of M on Y adjusted for X on the cluster or between-group level.
<code>se.a</code>	a positive numeric value indicating the standard error of a .
<code>se.b</code>	a positive numeric value indicating the standard error of b .
<code>cov.ab</code>	a positive numeric value indicating the covariance between a and b .
<code>cov.rand</code>	a positive numeric value indicating the covariance between the random slopes for a and b .
<code>se.cov.rand</code>	a positive numeric value indicating the standard error of the covariance between the random slopes for a and b .
<code>nrep</code>	an integer value indicating the number of Monte Carlo repetitions.
<code>alternative</code>	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
<code>seed</code>	a numeric value specifying the seed of the random number generator when using the Monte Carlo method.
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying
<code>write</code>	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

In statistical mediation analysis (MacKinnon & Tofighi, 2013), the indirect effect refers to the effect of the independent variable X on the outcome variable Y transmitted by the mediator variable M . The magnitude of the indirect effect ab is quantified by the product of the coefficient a (i.e., effect of X on M) and the coefficient b (i.e., effect of M on Y adjusted for X). However, mediation in the context of a 1-1-1 multilevel mediation model where variables X , M , and Y are measured at level 1, the coefficients a and b can vary across level-2 units (i.e., random slope). As a result, a and b may covary so that the estimate of the indirect effect is no longer simply the product of the coefficients $\hat{a}\hat{b}$, but $\hat{a}\hat{b} + \tau_{a,b}$, where $\tau_{a,b}$ (i.e., `cov.rand`) is the level-2 covariance between the random slopes a and b . The covariance term needs to be added to $\hat{a}\hat{b}$ only when random slopes are estimated for both a and b . Otherwise, the simple product is sufficient to quantify the indirect effect, and the `indirect` function can be used instead.

In practice, researchers are often interested in confidence limit estimation for the indirect effect. There are several methods for computing a confidence interval for the indirect effect in a single-level mediation models (see `indirect` function). The Monte Carlo (MC) method (MacKinnon et al., 2004) is a promising method in single-level mediation model which was also adapted to the multilevel mediation model (Bauer, Preacher & Gil, 2006). This method requires seven pieces of information available from the results of a multilevel mediation model:

- a** Coefficient a , i.e., average effect of X on M on the cluster or between-group level. In Mplus, Estimate of the random slope a under Means at the Between Level.
- b** Coefficient b , i.e., average effect of M on Y on the cluster or between-group level. In Mplus, Estimate of the random slope b under Means at the Between Level.
- se.a** Standard error of a . In Mplus, S.E. of the random slope a under Means at the Between Level.
- se.b** Standard error of b . In Mplus, S.E. of the random slope b under Means at the Between Level.
- cov.ab** Covariance between a and b . In Mplus, the estimated covariance matrix for the parameter estimates (i.e., asymptotic covariance matrix) need to be requested by specifying TECH3 along with TECH1 in the OUTPUT section. In the TECHNICAL 1 OUTPUT under PARAMETER SPECIFICATION FOR BETWEEN, the numbers of the parameter for the coefficients a and b need to be identified under ALPHA to look up `cov.av` in the corresponding row and column in the TECHNICAL 3 OUTPUT under ESTIMATED COVARIANCE MATRIX FOR PARAMETER ESTIMATES.
- cov.rand** Covariance between the random slopes for a and b . In Mplus, Estimate of the covariance a WITH b at the Between Level.
- se.cov.rand** Standard error of the covariance between the random slopes for a and b . In Mplus, S.E. of the covariance a WITH b at the Between Level.

Note that all pieces of information except `cov.ab` can be looked up in the standard output of the multilevel mediation model. In order to specify `cov.ab`, the covariance matrix for the parameter estimates (i.e., asymptotic covariance matrix) is required. In practice, `cov.ab` will oftentimes be very small so that `cov.ab` may be set to 0 (i.e., default value) with negligible impact on the results.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	list with the input specified in <code>a</code> , <code>b</code> , <code>se.a</code> , <code>se.b</code> , <code>cov.ab</code> , <code>cov.rand</code> , and <code>se.cov.rand</code>
<code>args</code>	specification of function arguments
<code>result</code>	list with result tables, i.e., <code>ab</code> for the simulated <code>ab</code> values and <code>mc</code> for the estimate of the indirect effect and the confidence interval

Note

The function was adapted from the interactive web tool by Preacher and Selig (2010).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Bauer, D. J., Preacher, K. J., & Gil, K. M. (2006). Conceptualizing and testing random indirect effects and moderated Mediation in multilevel models: New procedures and recommendations. *Psychological Methods*, 11, 142-163. <https://doi.org/10.1037/1082-989X.11.2.142>
- Kenny, D. A., Korchmaros, J. D., & Bolger, N. (2003). Lower level Mediation in multilevel models. *Psychological Methods*, 8, 115-128. <https://doi.org/10.1037/1082-989x.8.2.115>
- MacKinnon, D. P., Lockwood, C. M., & Williams, J. (2004). Confidence limits for the indirect effect: Distribution of the product and resampling methods. *Multivariate Behavioral Research*, 39, 99-128. https://doi.org/10.1207/s15327906mbr3901_4
- MacKinnon, D. P., & Tofighi, D. (2013). Statistical mediation analysis. In J. A. Schinka, W. F. Velicer, & I. B. Weiner (Eds.), *Handbook of psychology: Research methods in psychology* (pp. 717-735). John Wiley & Sons, Inc..
- Preacher, K. J., & Selig, J. P. (2010). *Monte Carlo method for assessing multilevel Mediation: An interactive tool for creating confidence intervals for indirect effects in 1-1-1 multilevel models* [Computer software]. Available from <http://quantpsy.org/>.

See Also

[indirect](#)

Examples

```
# Example 1: Confidence Interval for the Indirect Effect
multilevel.indirect(a = 0.25, b = 0.20, se.a = 0.11, se.b = 0.13,
                   cov.ab = 0.01, cov.rand = 0.40, se.cov.rand = 0.02)

# Example 2: Save results of the Monte Carlo method
ab <- multilevel.indirect(a = 0.25, b = 0.20, se.a = 0.11, se.b = 0.13,
                        cov.ab = 0.01, cov.rand = 0.40, se.cov.rand = 0.02,
                        output = FALSE)$result$ab

# Histogram of the distribution of the indirect effect
hist(ab)

## Not run:
# Example 3: Write results into a text file
multilevel.indirect(a = 0.25, b = 0.20, se.a = 0.11, se.b = 0.13,
                   cov.ab = 0.01, cov.rand = 0.40, se.cov.rand = 0.02,
                   write = "ML-Indirect.txt")

## End(Not run)
```

Description

This function evaluates configural, metric, and scalar cross-level measurement invariance using multilevel confirmatory factor analysis with continuous indicators by calling the `cfa` function in the R package **lavaan**.

Usage

```
multilevel.invar(data, ..., cluster, model = NULL, rescov = NULL,
  invar = c("config", "metric", "scalar"), fix.resid = NULL,
  ident = c("marker", "var", "effect"),
  estimator = c("ML", "MLR"), optim.method = c("nlminb", "em"),
  missing = c("listwise", "fiml"),
  print = c("all", "summary", "coverage", "descript", "fit",
    "est", "modind", "resid"),
  print.fit = c("all", "standard", "scaled", "robust"),
  mod.minval = 6.63, resid.minval = 0.1, digits = 3, p.digits = 3,
  as.na = NULL, write = NULL, append = TRUE, check = TRUE,
  output = TRUE)
```

Arguments

- | | |
|---------|---|
| data | a data frame. If <code>model</code> is <code>NULL</code> , multilevel confirmatory factor analysis based on a measurement model with one factor at the Within and Between level comprising all variables in the data frame is conducted to evaluate cross-level measurement invariance. Note that the cluster variable specified in <code>cluster</code> is excluded from data when specifying the argument <code>cluster</code> using the variable name of the cluster variable. If <code>model</code> is specified, the data frame needs to contain all variables used in the <code>model</code> argument. |
| ... | an expression indicating the variable names in data, e.g., <code>multilevel.invar(dat, x1, x2, x3, cluster = "cluster")</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function. |
| cluster | either a character string indicating the variable name of the cluster variable in data, or a vector representing the nested grouping structure (i.e., group or cluster variable). |
| model | a character vector specifying the same factor structure with one factor at the Within and Between Level, or a list of character vectors for specifying the same measurement model with more than one factor at the Within and Between Level, e.g., <code>model = c("x1", "x2", "x3", "x4")</code> for specifying a measurement model with one factor labeled <code>wf</code> at the Within level and a measurement model with one factor labeled <code>bf</code> at the Between level each comprising four indicators, or <code>model = list(factor1 = c("x1", "x2", "x3", "x4"), factor2 = c("x5", "x6", "x7", "x8"))</code> for specifying a measurement model with two latent factors labeled <code>wfactor1</code> and <code>wfactor2</code> at the Within level and a measurement model with two latent factors labeled <code>bfactor1</code> and <code>bfactor2</code> at the Between level each comprising four indicators. Note that the name of each list element is used to label factors, where prefixes <code>w</code> and <code>b</code> are added the labels to distinguish factor labels at the Within and Between level, i.e., all list elements need to be |

named, otherwise factors are labeled with "wf1", "wf2", "wf3" for labels at the Within level and "bf1", "bf2", "bf3" for labels at the Between level and so on.

rescov	a character vector or a list of character vectors for specifying residual covariances at the Within level, e.g. <code>rescov = c("x1", "x2")</code> for specifying a residual covariance between indicators x1 and x2 at the Within level or <code>rescov = list(c("x1", "x2"), c("x3", "x4"))</code> for specifying residual covariances between indicators x1 and x2, and indicators x3 and x4 at the Within level. Note that residual covariances at the Between level can only be specified by using the arguments <code>model.w</code> , <code>model.b</code> , and <code>model.b</code> .
invar	a character string indicating the level of measurement invariance to be evaluated, i.e., <code>config</code> to evaluate configural measurement invariance (i.e., same factor structure across levels), <code>metric</code> (default) to evaluate configural and metric measurement invariance (i.e., equal factor loadings across level), and <code>scalar</code> to evaluate configural, metric and scalar measurement invariance (i.e., all residual variances at the Between level equal zero).
fix.resid	a character vector for specifying residual variances to be fixed at 0 at the Between level for the configural and metric invariance model, e.g., <code>fix.resid = c("x1", "x3")</code> to fix residual variances of indicators x1 and x2 at the Between level at 0. Note that it is also possible to specify <code>fix.resid = "all"</code> which fixes all residual variances at the Between level at 0 in line with the strong factorial measurement invariance assumption across cluster.
ident	a character string indicating the method used for identifying and scaling latent variables, i.e., <code>"marker"</code> for the marker variable method fixing the first factor loading of each latent variable to 1, <code>"var"</code> for the fixed variance method fixing the variance of each latent variable to 1, or <code>"effect"</code> for the effects-coding method using equality constraints so that the average of the factor loading for each latent variable equals 1.
estimator	a character string indicating the estimator to be used: <code>"ML"</code> for maximum likelihood with conventional standard errors and <code>"MLR"</code> (default) for maximum likelihood with Huber-White robust standard errors and a scaled test statistic that is asymptotically equal to the Yuan-Bentler test statistic. Note that by default, full information maximum likelihood (FIML) method is used to deal with missing data when using <code>"ML"</code> (<code>missing = "fiml"</code>), whereas incomplete cases are removed listwise (i.e., <code>missing = "listwise"</code>) when using <code>"MLR"</code> .
optim.method	a character string indicating the optimizer, i.e., <code>"nlminb"</code> (default) for the unconstrained and bounds-constrained quasi-Newton method optimizer and <code>"em"</code> for the Expectation Maximization (EM) algorithm.
missing	a character string indicating how to deal with missing data, i.e., <code>"listwise"</code> (default) for listwise deletion or <code>"fiml"</code> for full information maximum likelihood (FIML) method. Note that FIML method is only available when <code>estimator = "ML"</code> , that it takes longer to estimate the model using FIML, and that FIML is prone to convergence issues which might be resolved by switching to listwise deletion.
print	a character string or character vector indicating which results to show on the console, i.e. <code>"all"</code> for all results, <code>"summary"</code> for a summary of the specification of the estimation method and missing data handling in lavaan, <code>"coverage"</code>

	for the variance-covariance coverage of the data, "descript" for descriptive statistics, "fit" for model fit and model comparison, "est" for parameter estimates, and "modind" for modification indices. By default, a summary of the specification and model fit and model comparison are printed.
print.fit	a character string or character vector indicating which version of the CFI, TLI, and RMSEA to show on the console, i.e., "all" for all versions of the CFI, TLI, and RMSEA, "standard" (default when estimator = "ML") for fit indices without any non-normality correction, "scaled" for population-corrected robust fit indices with ad hoc non-normality correction, and robust (default when estimator = "MLR") for sample-corrected robust fit indices based on formula provided by Li and Bentler (2006) and Brosseau-Liard and Savalei (2014).
mod.minval	numeric value to filter modification indices and only show modifications with a modification index value equal or higher than this minimum value. By default, modification indices equal or higher 6.63 are printed. Note that a modification index value of 6.63 is equivalent to a significance level of $\alpha = .01$.
resid.minval	numeric value indicating the minimum absolute residual correlation coefficients and standardized means to highlight in boldface. By default, absolute residual correlation coefficients and standardized means equal or higher 0.1 are highlighted. Note that highlighting can be disabled by setting the minimum value to 1.
digits	an integer value indicating the number of decimal places to be used for displaying results. Note that information criteria and chi-square test statistic is printed with digits minus 1 decimal places.
p.digits	an integer value indicating the number of decimal places to be used for displaying the <i>p</i> -value.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that as.na() function is only applied to data but not to cluster.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification, convergence and model identification is checked.
output	logical: if TRUE (default), output is shown.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	data frame specified in data

args	specification of function arguments
model	list with specified model for the configural, metric, and scalar invariance model
model.fit	list with fitted lavaan object of the configural, metric, and scalar invariance model
check	list with the results of the convergence and model identification check for the configural, metric, and scalar invariance model
result	list with result tables, i.e., summary for the summary of the specification of the estimation method and missing data handling in lavaan, coverage for the variance-covariance coverage of the data, <code>descript</code> for descriptive statistics, <code>fit</code> for a list with model fit based on standard, scaled, and robust fit indices, <code>est</code> for a list with parameter estimates for the configural, metric, and scalar invariance model, and <code>modind</code> for the list with modification indices for the configural, metric, and scalar invariance model

Note

The function uses the functions `lavTestLRT` provided in the R package **lavaan** by Yves Rosseel (2012).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Rosseel, Y. (2012). lavaan: An R Package for Structural Equation Modeling. *Journal of Statistical Software*, 48, 1-36. <https://doi.org/10.18637/jss.v048.i02>

See Also

[multilevel.cfa](#), [multilevel.fit](#), [multilevel.omega](#), [multilevel.cor](#), [multilevel.descript](#)

Examples

```
## Not run:

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

#-----
# Cluster variable specification

# Example 1a: Specification using the argument '...'
multilevel.invar(Demo.twolevel, y1:y4, cluster = "cluster")

# Example 1b: Alternative specification with cluster variable 'cluster' in 'data'
multilevel.invar(Demo.twolevel[, c("y1", "y2", "y3", "y4", "cluster")], cluster = "cluster")

# Example 1b: Alternative specification with cluster variable 'cluster' not in 'data'
multilevel.invar(Demo.twolevel[, c("y1", "y2", "y3", "y4")], cluster = Demo.twolevel$cluster)
```

```

#-----
# Model specification using 'data' for a one-factor model

#.....
# Level of measurement invariance

# Example 2a: Configural invariance
multilevel.invar(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", invar = "config")

# Example 2b: Metric invariance
multilevel.invar(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", invar = "metric")

# Example 2c: Scalar invariance
multilevel.invar(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", invar = "scalar")

#.....
# Residual covariance at the Within level and residual variance at the Between level

# Example 3a: Residual covariance between "y3" and "y4" at the Within level
multilevel.invar(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster",
                 rescov = c("y3", "y4"))

# Example 3b: Residual variances of 'y1' at the Between level fixed at 0
multilevel.invar(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", fix.resid = "y1")

#.....
# Example 4: Print all results
multilevel.invar(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", print = "all")

#.....
# Example 5: lavaan model and summary of the estimated model
mod <- multilevel.invar(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", output = FALSE)

# lavaan syntax of the metric invariance model
mod$model$metric

# Fitted lavaan object of the metric invariance model
lavaan::summary(mod$model.fit$metric, standardized = TRUE, fit.measures = TRUE)

#-----
# Model specification using 'model' for one or multiple factor model

# Example 6a: One-factor model
multilevel.invar(Demo.twolevel, cluster = "cluster", model = c("y1", "y2", "y3", "y4"))

# Example 6b: Two-factor model
multilevel.invar(Demo.twolevel, cluster = "cluster",
                 model = list(c("y1", "y2", "y3"), c("y4", "y5", "y6")))

#-----
# Write results

```

```
# Example 7a: Write Results into a Excel file
multilevel.invar(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", print = "all",
  write = "Multilevel_Invariance.txt")

# Example 7b: Write Results into a Excel file
multilevel.invar(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", print = "all",
  write = "Multilevel_Invariance.xlsx")

## End(Not run)
```

multilevel.omega

Multilevel Composite Reliability

Description

This function computes point estimate and Monte Carlo confidence interval for the multilevel composite reliability defined by Lai (2021) for a within-cluster construct, shared cluster-level construct, and configural cluster construct by calling the `cfa` function in the R package **lavaan**.

Usage

```
multilevel.omega(data, ..., cluster, rescov = NULL,
  const = c("within", "shared", "config"),
  fix.resid = NULL, optim.method = c("nlminb", "em"),
  missing = c("listwise", "fiml"), nrep = 100000, seed = NULL,
  conf.level = 0.95, print = c("all", "omega", "item"),
  digits = 2, as.na = NULL, write = NULL, append = TRUE,
  check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a data frame. Multilevel confirmatory factor analysis based on a measurement model with one factor at the Within level and one factor at the Between level comprising all variables in the data frame is conducted. Note that the cluster variable specified in <code>cluster</code> is excluded from data when specifying the argument <code>cluster</code> using the variable name of the cluster variable.
<code>...</code>	an expression indicating the variable names in data, e.g., <code>multilevel.omega(dat, x1, x2, x3, cluster = "cluster")</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>cluster</code>	either a character string indicating the variable name of the cluster variable in data, or a vector representing the nested grouping structure (i.e., group or cluster variable).
<code>rescov</code>	a character vector or a list of character vectors for specifying residual covariances at the Within level, e.g. <code>rescov = c("x1", "x2")</code> for specifying a residual covariance between indicators <code>x1</code> and <code>x2</code> at the Within level or <code>rescov = list(c("x1", "x2"), c("x3", "x4"))</code> for specifying residual covariances between indicators <code>x1</code> and <code>x2</code> , and indicators <code>x3</code> and <code>x4</code> at the Within level. Note

	that residual covariances at the Between level cannot be specified using this function.
const	a character string indicating the type of construct(s), i.e., "within" for within-cluster constructs, "shared" for shared cluster-level constructs, and "config" (default) for configural cluster constructs.
fix.resid	a character vector for specifying residual variances to be fixed at 0 at the Between level, e.g., <code>fix.resid = c("x1", "x3")</code> to fix residual variances of indicators x1 and x2 at the Between level at 0. Note that it is also possible to specify <code>fix.resid = "all"</code> which fixes all residual variances at the Between level at 0 in line with the strong factorial measurement invariance assumption across cluster.
optim.method	a character string indicating the optimizer, i.e., "nlminb" (default) for the unconstrained and bounds-constrained quasi-Newton method optimizer and "em" for the Expectation Maximization (EM) algorithm.
missing	a character string indicating how to deal with missing data, i.e., "listwise" for listwise deletion or "fiml" (default) for full information maximum likelihood (FIML) method.
nrep	an integer value indicating the number of Monte Carlo repetitions for computing confidence intervals.
seed	a numeric value specifying the seed of the random number generator for computing the Monte Carlo confidence interval.
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
print	a character vector indicating which results to show, i.e. "all" (default), for all results "omega" for omega, and "item" for item statistics.
digits	an integer value indicating the number of decimal places to be used for displaying results. Note that loglikelihood, information criteria and chi-square test statistic is printed with digits minus 1 decimal places.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis. Note that <code>as.na()</code> function is only applied to data but not to cluster.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification, convergence and model identification is checked.
output	logical: if TRUE (default), output is shown.

Value

call	function call
type	type of analysis

data	data frame specified in data including the group variable specified in cluster
args	specification of function arguments
model	specified model
model.fit	fitted lavaan object (mod.fit)
check	results of the convergence and model identification check
result	list with result tables, i.e., omega for the coefficient omega including Monte Carlo confidence interval and itemstat for descriptive statistics

Note

The function uses the functions `lavInspect`, `lavTech`, and `lavNames`, provided in the R package **lavaan** by Yves Rosseel (2012). The internal function `.internal.mvrnorm` is a copy of the `mvrnorm` function in the package **MASS** by Venables and Ripley (2002).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Lai, M. H. C. (2021). Composite reliability of multilevel data: It's about observed scores and construct meanings. *Psychological Methods*, 26(1), 90–102. <https://doi.org/10.1037/met0000287>

Rosseel, Y. (2012). lavaan: An R Package for Structural Equation Modeling. *Journal of Statistical Software*, 48, 1-36. <https://doi.org/10.18637/jss.v048.i02>

Venables, W. N., Ripley, B. D. (2002). *Modern Applied Statistics with S* (4th ed.). Springer. <https://www.stats.ox.ac.uk/pub/MASS4>

See Also

[item.omega](#), [multilevel.cfa](#), [multilevel.fit](#), [multilevel.invar](#), [multilevel.cor](#), [multilevel.descript](#)

Examples

```
## Not run:

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

#-----
# Cluster variable specification

# Example 1a: Specification using the argument '...'
multilevel.omega(Demo.twolevel, y1:y4, cluster = "cluster")

# Example 1b: Alternative specification with cluster variable 'cluster' in 'data'
multilevel.omega(Demo.twolevel[, c("y1", "y2", "y3", "y4", "cluster")], cluster = "cluster")

# Example 1b: Alternative specification with cluster variable 'cluster' not in 'data'
multilevel.omega(Demo.twolevel[, c("y1", "y2", "y3", "y4")], cluster = Demo.twolevel$cluster)
```

```

#-----
# Type of construct

# Example 2a: Within-Cluster Construct
multilevel.omega(Demo.twolevel[, c("y1", "y2", "y3", "y4")],
                 cluster = Demo.twolevel$cluster, const = "within")

# Example 2b: Shared Cluster-Level Construct
multilevel.omega(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", const = "shared")

# Example 2c: Configural Construct
multilevel.omega(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", const = "config")

#-----
# Residual covariance at the Within level and residual variance at the Between level

# Example 3a: Residual covariance between "y4" and "y5" at the Within level
multilevel.omega(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", const = "config",
                 rescov = c("y3", "y4"))

# Example 3b: Residual variances of 'y1' at the Between level fixed at 0
multilevel.omega(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster", const = "config",
                 fix.resid = c("y1", "y2"), digits = 3)

#-----
# Write results

# Example 4a: Write results into a text file
multilevel.omega(Demo.twolevel[, c("y1", "y2", "y3", "y4")],
                 cluster = Demo.twolevel$cluster, write = "Multilevel_Omega.txt")

# Example 4b: Write results into a Excel file
multilevel.omega(Demo.twolevel, y1, y2, y3, y4, cluster = "cluster",
                 write = "Multilevel_Omega.xlsx")

## End(Not run)

```

multilevel.r2

R-Squared Measures for Multilevel and Linear Mixed Effects Models

Description

This function computes R-squared measures by Raudenbush and Bryk (2002), Snijders and Bosker (1994), Nakagawa and Schielzeth (2013) as extended by Johnson (2014), and Rights and Sterba (2019) for multilevel and linear mixed effects models estimated by using the `lmer()` function in the package **lme4** or `lme()` function in the package **nlme**.

Usage

```
multilevel.r2(model, print = c("all", "RB", "SB", "NS", "RS"), digits = 3,
```

```

plot = FALSE, gray = FALSE, start = 0.15, end = 0.85,
color = c("#D55E00", "#0072B2", "#CC79A7", "#009E73", "#E69F00"),
filename = NULL, width = NA, height = NA,
units = c("in", "cm", "mm", "px"), dpi = 600,
write = NULL, append = TRUE, check = TRUE, output = TRUE)

```

Arguments

model	a fitted model of class "lmerMod" from the lme4 package or "lme" from the nlme package.
print	a character vector indicating which R-squared measures to be printed on the console, i.e., RB for measures from Raudenbush and Bryk (2002), SB for measures from Snijders and Bosker (1994), NS for measures from Nakagawa and Schielzeth (2013) as extended by Johnson (2014), and RS for measures from Rights and Sterba (2019). The default setting is print = "RS".
digits	an integer value indicating the number of decimal places to be used.
plot	logical: if TRUE, bar chart showing the decomposition of scaled total, within-cluster, and between-cluster outcome variance into five (total), three (within-cluster), and two (between-cluster) proportions is drawn. Note that the ggplot2 package is required to draw the bar chart.
gray	logical: if TRUE, graphical parameter to draw the bar chart in gray scale.
start	a numeric value between 0 and 1, graphical parameter to specify the gray value at the low end of the palette.
end	a numeric value between 0 and 1, graphical parameter to specify the gray value at the high end of the palette.
color	a character vector, graphical parameter indicating the color of bars in the bar chart in the following order: Fixed slopes (Within), Fixed slopes (Between), Slope variation (Within), Intercept variation (Between), and Residual (Within). By default, colors from the colorblind-friendly palettes are used.
filename	a character string indicating the filename argument including the file extension in the ggsave function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the file argument. Note that plots can only be saved when plot = TRUE and print = "RS".
width	a numeric value indicating the width argument (default is the size of the current graphics device) in the ggsave function.
height	a numeric value indicating the height argument (default is the size of the current graphics device) in the ggsave function.
units	a character string indicating the units argument (default is in) in the ggsave function.
dpi	a numeric value indicating the dpi argument (default is 600) in the ggsave function.
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.

append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

A number of R-squared measures for multilevel and linear mixed effects models have been developed in the methodological literature (see Rights & Sterba, 2018). Based on these measures, following measures were implemented in the current function:

Raudenbush and Bryk (2002) R-squared measures by Raudenbush and Bryk (2002) are based on the proportional reduction of unexplained variance when predictors are added. More specifically, variance estimates from the baseline/null model (i.e., $\sigma_{e|b}^2$ and $\sigma_{u0|b}^2$) and variance estimates from the model including predictors (i.e., $\sigma_{e|m}^2$ and $\sigma_{u0|m}^2$) are used to compute the proportional reduction in variance between baseline/null model and the complete model by:

$$R_1^2(RB) = \frac{\sigma_{e|b}^2 - \sigma_{e|m}^2}{\sigma_{e|b}^2}$$

for the proportional reduction at level-1 (within-cluster) and by:

$$R_2^2(RB) = \frac{\sigma_{u0|b}^2 - \sigma_{u0|m}^2}{\sigma_{u0|b}^2}$$

for the proportional reduction at level-2 (between-cluster), where $|b$ and $|m$ represent the baseline and full models, respectively (Hox et al., 2018; Roberts et al., 2010).

A major disadvantage of these measures is that adding predictors can increase rather than decrease some of the variance components and it is even possible to obtain negative values for R^2 with these formulas (Snijders & Bosker, 2012). According to Snijders and Bosker (1994) this can occur because the between-group variance is a function of both level-1 and level-2 variance:

$$var(\bar{Y}_j) = \sigma_{u0}^2 + \frac{\sigma_e^2}{n_j}$$

Hence, adding a predictor (e.g., cluster-mean centered predictor) that explains proportion of the within-group variance will decrease the estimate of σ_e^2 and increase the estimate σ_{u0}^2 if this predictor does not explain a proportion of the between-group variance to balance out the decrease in σ_e^2 (LaHuis et al., 2014). Negative estimates for R^2 can also simply occur due to chance fluctuation in sample estimates from the two models.

Another disadvantage of these measures is that $R_2^2(RB)$ for the explained variance at level-2 has been shown to perform poorly in simulation studies even with $j = 200$ clusters with group cluster size of $n_j = 50$ (LaHuis et al., 2014; Rights & Sterba, 2019).

Moreover, when there is missing data in the level-1 predictors, it is possible that sample sizes for the baseline and complete models differ.

Finally, it should be noted that R-squared measures by Raudenbush and Bryk (2002) are appropriate for random intercept models, but not for random intercept and slope models. For

random slope models, Snijders and Bosker (2012) suggested to re-estimate the model as random intercept models with the same predictors while omitting the random slopes to compute the R-squared measures. However, the simulation study by LaHuis (2014) suggested that the R-squared measures showed an acceptable performance when there was little slope variance, but did not perform well in the presence of higher levels of slope variance.

Snijders and Bosker (1994) R-squared measures by Snijders and Bosker (1994) are based on the proportional reduction of mean squared prediction error and is computed using the formula:

$$R_1^2(SB) = \frac{\hat{\sigma}_{e|m}^2 + \hat{\sigma}_{u0|m}^2}{\hat{\sigma}_{e|b}^2 + \hat{\sigma}_{u0|b}^2}$$

for computing the proportional reduction of error at level-1 representing the total amount of explained variance and using the formula:

$$R_2^2(SB) = \frac{\hat{\sigma}_{e|m}^2/n_j + \hat{\sigma}_{u0|m}^2}{\hat{\sigma}_{e|b}^2/n_j + \hat{\sigma}_{u0|b}^2}$$

for computing the proportional reduction of error at level-2 by dividing the $\hat{\sigma}_e^2$ by the group cluster size n_j or by the average cluster size for unbalanced data (Roberts et al., 2010). Note that the function uses the harmonic mean of the group sizes as recommended by Snijders and Bosker (1994). The population values of R^2 based on these measures cannot be negative because the interplay of level-1 and level-2 variance components is considered. However, sample estimates of R^2 can be negative either due to chance fluctuation when sample sizes are small or due to model misspecification (Snijders and Bosker, 2012).

When there is missing data in the level-1 predictors, it is possible that sample sizes for the baseline and complete models differ.

Similar to the R-squared measures by Raudenbush and Bryk (2002), the measures by Snijders and Bosker (1994) are appropriate for random intercept models, but not for random intercept and slope models. Accordingly, for random slope models, Snijders and Bosker (2012) suggested to re-estimate the model as random intercept models with the same predictors while omitting the random slopes to compute the R-squared measures. The simulation study by LaHuis et al. (2014) revealed that the R-squared measures showed an acceptable performance, but it should be noted that $R_2^2(SB)$ the explained variance at level-2 was not investigated in their study.

Nakagawa and Schielzeth (2013) R-squared measures by Nakagawa and Schielzeth (2013) are based on partitioning model-implied variance from a single fitted model and uses the variance of predicted values of $\text{var}(\hat{Y}_{ij})$ to form both the outcome variance in the denominator and the explained variance in the numerator of the formulas:

$$R_m^2(NS) = \frac{\text{var}(\hat{Y}_{ij})}{\text{var}(\hat{Y}_{ij}) + \sigma_{u0}^2 + \sigma_e^2}$$

for marginal total $R_m^2(NS)$ and:

$$R_c^2(NS) = \frac{\text{var}(\hat{Y}_{ij}) + \sigma_{u0}^2}{\text{var}(\hat{Y}_{ij}) + \sigma_{u0}^2 + \sigma_e^2}$$

for conditional total $R_c^2(NS)$. In the former formula R^2 predicted scores are marginalized across random effects to indicate the variance explained by fixed effects and in the latter formula R^2 predicted scores are conditioned on random effects to indicate the variance explained by fixed and random effects (Rights and Sterba, 2019).

The advantage of these measures is that they can never become negative and that they can also be extended to generalized linear mixed effects models (GLMM) when outcome variables are not continuous (e.g., binary outcome variables). Note that currently the function does not provide R^2 measures for GLMMs, but these measures can be obtained using the `r.squaredGLMM()` function in the **MuMIn** package.

A disadvantage is that these measures do not allow random slopes and are restricted to the simplest random effect structure (i.e., random intercept model). In other words, these measures do not fully reflect the structure of the fitted model when using random intercept and slope models. However, Johnson (2014) extended these measures to allow random slope by taking into account the contribution of random slopes, intercept-slope covariances, and the covariance matrix of random slope to the variance in Y_{ij} . As a result, R-squared measures by Nakagawa and Schielzeth (2013) as extended by Johnson (2014) can be used for both random intercept, and random intercept and slope models.

The major criticism of the R-squared measures by Nakagawa and Schielzeth (2013) as extended by Johnson (2014) is that these measures do not decompose outcome variance into each of total, within-cluster, and between-cluster variance which precludes from computing level-specific R^2 measures. In addition, these measures do not distinguish variance attributable to level-1 versus level-2 predictors via fixed effects, and they also do not distinguish between random intercept and random slope variation (Rights and Sterba, 2019).

Rights and Sterba (2019) R-squared measures by Rights and Sterba (2019) provide an integrative framework of R-squared measures for multilevel and linear mixed effects models with random intercepts and/or slopes. Their measures are also based on partitioning model implied variance from a single fitted model, but they provide a full partitioning of the total outcome variance to one of five specific sources:

- variance attributable to level-1 predictors via fixed slopes (shorthand: variance attributable to f1)
- variance attributable to level-2 predictors via fixed slopes (shorthand: variance attributable to f2)
- variance attributable to level-1 predictors via random slope variation/ covariation (shorthand: variance attributable to v)
- variance attributable to cluster-specific outcome means via random intercept variation (shorthand: variance attributable to m)
- variance attributable to level-1 residuals

R^2 measures are based on the outcome variance of interest (total, within-cluster, or between-cluster) in the denominator, and the source contributing to explained variance in the numerator:

Total R^2 measures incorporate both within-cluster and between cluster variance in the denominator and quantify variance explained in an omnibus sense:

- $R_t^{2(f_1)}$: Proportion of total outcome variance explained by level-1 predictors via fixed slopes.
- $R_t^{2(f_2)}$: Proportion of total outcome variance explained by level-2 predictors via fixed slopes.

- $R_t^{2(f)}$: Proportion of total outcome variance explained by all predictors via fixed slopes.
- $R_t^{2(v)}$: Proportion of total outcome variance explained by level-1 predictors via random slope variation/covariation.
- $R_t^{2(m)}$: Proportion of total outcome variance explained by cluster-specific outcome means via random intercept variation.
- $R_t^{2(fv)}$: Proportion of total outcome variance explained by predictors via fixed slopes and random slope variation/covariation.
- $R_t^{2(fvm)}$: Proportion of total outcome variance explained by predictors via fixed slopes and random slope variation/covariation and by cluster-specific outcome means via random intercept variation.

Within-Cluster R^2 measures incorporate only within-cluster variance in the denominator and indicate the degree to which within-cluster variance can be explained by a given model:

- $R_w^{2(f_1)}$: Proportion of within-cluster outcome variance explained by level-1 predictors via fixed slopes.
- $R_w^{2(v)}$: Proportion of within-cluster outcome variance explained by level-1 predictors via random slope variation/covariation.
- $R_w^{2(f_1v)}$: Proportion of within-cluster outcome variance explained by level-1 predictors via fixed slopes and random slope variation/covariation.

Between-Cluster R^2 measures incorporate only between-cluster variance in the denominator and indicate the degree to which between-cluster variance can be explained by a given model:

- $R_b^{2(f_2)}$: Proportion of between-cluster outcome variance explained by level-2 predictors via fixed slopes.
- $R_b^{2(m)}$: Proportion of between-cluster outcome variance explained by cluster-specific outcome means via random intercept variation.

The decomposition of the total outcome variance can be visualized in a bar chart by specifying `plot = TRUE`. The first column of the bar chart decomposes scaled total variance into five distinct proportions (i.e., $R_t^{2(f_1)}$, $R_t^{2(f_2)}$, $R_t^{2(f)}$, $R_t^{2(v)}$, $R_t^{2(m)}$, $R_t^{2(fv)}$, and $R_t^{2(fvm)}$), the second column decomposes scaled within-cluster variance into three distinct proportions (i.e., $R_w^{2(f_1)}$, $R_w^{2(v)}$, and $R_w^{2(f_1v)}$), and the third column decomposes scaled between-cluster variance into two distinct proportions (i.e., $R_b^{2(f_2)}$, $R_b^{2(m)}$).

Note that the function assumes that all level-1 predictors are centered within cluster (i.e., group-mean or cluster-mean centering) as has been widely recommended (e.g., Enders & Tofghi, D., 2007; Rights et al., 2019). In fact, it does not matter whether a lower-level predictor is merely a control variable, or is quantitative or categorical (Yaremych et al., 2021), cluster-mean centering should always be used for lower-level predictors to obtain an orthogonal between-within partitioning of a lower-level predictor's variance that directly parallels what happens to a level-1 outcome (Hoffman & Walters, 2022). In the absence of cluster-mean-centering, however, the function provides total R^2 measures, but does not provide any within-cluster or between-cluster R^2 measures.

By default, the function only computes R-squared measures by Rights and Sterba (2019) because the other R-squared measures reflect the same population quantity provided by Rights and Sterba

(2019). That is, R-squared measures $R_1^2(RB)$ and $R_2^2(RB)$ by Raudenbush and Bryk (2002) are equivalent to $R_w^{2(f_1v)}$ and $R_b^{2(f_2)}$, R-squared measures $R_1^2(SB)$ and $R_2^2(SB)$ are equivalent to $R_t^{2(f)}$ and $R_b^{2(f_2)}$, and R-squared measures $R_m^2(NS)$ and $R_c^2(NS)$ by Nakagawa and Schielzeth (2013) as extended by Johnson (2014) are equivalent to $R_t^{2(f)}$ and $R_t^{2(fvm)}$ (see Rights and Sterba, Table 3).

Note that none of these measures provide an R^2 for the random slope variance explained by cross-level interactions, a quantity that is frequently of interest (Hoffman & Walters, 2022).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	matrix or data frame specified in data
<code>plot</code>	ggplot2 object for plotting the results
<code>args</code>	specification of function arguments
<code>result</code>	list with result tables, i.e., <code>rb</code> for the R2 measures by Raudenbush and Bryk (2002), <code>sb</code> for the R2 measures by Snijders and Bosker (1994), <code>ns</code> for the R2 measures by Nakagawa and Schielzeth (2013), and <code>rs</code> for the R2 measures by Rights and Sterba (2019)

Note

This function is based on the `multilevelR2()` function from the **mitml** package by Simon Grund, Alexander Robitzsch and Oliver Luedtke (2021), and a copy of the function `r2mlm` in the **r2mlm** package by Mairead Shaw, Jason Rights, Sonya Sterba, and Jessica Flake.

Author(s)

Simon Grund, Alexander Robitzsch, Oliver Luedtk, Mairead Shaw, Jason D. Rights, Sonya K. Sterba, Jessica K. Flake, and Takuya Yanagida

References

- Enders, C. K., & Tofighi, D. (2007). Centering predictor variables in cross-sectional multilevel models: A new look at an old issue. *Psychological Methods*, 12, 121-138. <https://doi.org/10.1037/1082-989X.12.2.121>
- Hoffmann, L., & Walter, W. R. (2022). Catching up on multilevel modeling. *Annual Review of Psychology*, 73, 629-658. <https://doi.org/10.1146/annurev-psych-020821-103525>
- Hox, J., Moerbeek, M., & van de Schoot, R. (2018). *Multilevel Analysis: Techniques and Applications* (3rd ed.) Routledge.
- Johnson, P. C. D. (2014). Extension of Nakagawa & Schielzeth's R2 GLMM to random slopes models. *Methods in Ecology and Evolution*, 5(9), 944-946. <https://doi.org/10.1111/2041-210X.12225>
- LaHuis, D. M., Hartman, M. J., Hakoyama, S., & Clark, P. C. (2014). Explained variance measures for multilevel models. *Organizational Research Methods*, 17, 433-451. <https://doi.org/10.1177/1094428114541701>

- Nakagawa, S., & Schielzeth, H. (2013). A general and simple method for obtaining R² from generalized linear mixed-effects models. *Methods in Ecology and Evolution*, 4(2), 133-142. <https://doi.org/10.1111/j.2041-210x.2012.00261.x>
- Raudenbush, S. W., & Bryk, A. S., (2002). *Hierarchical linear models: Applications and data analysis methods*. Sage.
- Rights, J. D., Preacher, K. J., & Cole, D. A. (2020). The danger of conflating level-specific effects of control variables when primary interest lies in level-2 effects. *British Journal of Mathematical and Statistical Psychology*, 73(Suppl 1), 194-211. <https://doi.org/10.1111/bmsp.12194>
- Rights, J. D., & Sterba, S. K. (2019). Quantifying explained variance in multilevel models: An integrative framework for defining R-squared measures. *Psychological Methods*, 24, 309-338. <https://doi.org/10.1037/met0000184>
- Roberts, K. J., Monaco, J. P., Stovall, H., & Foster, V. (2011). Explained variance in multilevel models (pp. 219-230). In J. J. Hox & J. K. Roberts (Eds.), *Handbook of Advanced Multilevel Analysis*. Routledge.
- Snijders, T. A. B., & Bosker, R. (1994). Modeled variance in two-level models. *Sociological methods and research*, 22, 342-363. <https://doi.org/10.1177/0049124194022003004>
- Snijders, T. A. B., & Bosker, R. J. (2012). *Multilevel analysis: An introduction to basic and advanced multilevel modeling* (2nd ed.). Sage.
- Yaremych, H. E., Preacher, K. J., & Hedeker, D. (2021). Centering categorical predictors in multilevel models: Best practices and interpretation. *Psychological Methods*. Advance online publication. <https://doi.org/10.1037/met0000434>

See Also

[multilevel.cor](#), [multilevel.descript](#), [multilevel.icc](#), [multilevel.indirect](#)

Examples

```
## Not run:

# Load misty, lme4, nlme, and ggplot2 package
misty::libraries(misty, lme4, nlme, ggplot2)

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

#-----

# Cluster mean centering, center() from the misty package
Demo.twolevel <- center(Demo.twolevel, x2, type = "CWC", cluster = "cluster")

# Compute group means, cluster.scores() from the misty package
Demo.twolevel <- cluster.scores(Demo.twolevel, x2, cluster = "cluster", name = "x2.b")

# Estimate multilevel model using the lme4 package
mod1a <- lmer(y1 ~ x2.c + x2.b + w1 + (1 + x2.c | cluster), data = Demo.twolevel,
              REML = FALSE, control = lmerControl(optimizer = "bobyqa"))
```

```

# Estimate multilevel model using the nlme package
mod1b <- lme(y1 ~ x2.c + x2.b + w1, random = ~ 1 + x2.c | cluster, data = Demo.twolevel,
            method = "ML")

#-----
# Example 1a: R-squared measures according to Rights and Sterba (2019)
multilevel.r2(mod1a)

# Example 1b: R-squared measures according to Rights and Sterba (2019)
multilevel.r2(mod1b)

# Example 1a: Write Results into a text file
multilevel.r2(mod1a, write = "ML-R2.txt")

#-----
# Example 2: Bar chart showing the decomposition of scaled total, within-cluster,
# and between-cluster outcome variance
multilevel.r2(mod1a, plot = TRUE)

# Bar chart in gray scale
multilevel.r2(mod1a, plot = TRUE, gray = TRUE)

# Save bar chart
multilevel.r2(mod1a, plot = TRUE, filename = "Proportion_of_Variance.png",
              dpi = 600, width = 5.5, height = 5.5)

#-----
# Example 3: Estimate multilevel model without random slopes
# Note. R-squared measures by Raudenbush and Bryk (2002), and Snijders and
# Bosker (2012) should be computed based on the random intercept model
mod2 <- lmer(y1 ~ x2.c + x2.b + w1 + (1 | cluster), data = Demo.twolevel,
            REML = FALSE, control = lmerControl(optimizer = "bobyqa"))

# Print all available R-squared measures
multilevel.r2(mod2, print = "all")

#-----
# Example 4: Draw bar chart manually
mod1a.r2 <- multilevel.r2(mod1a, output = FALSE)

# Prepare data frame for ggplot()
df <- data.frame(var = factor(rep(c("Total", "Within", "Between"), each = 5),
                             level = c("Total", "Within", "Between")),
                 part = factor(c("Fixed Slopes (Within)", "Fixed Slopes (Between)",
                                "Slope Variation (Within)", "Intercept Variation (Between)",
                                "Residual (Within)"),
                             level = c("Residual (Within)", "Intercept Variation (Between)",
                                "Slope Variation (Within)", "Fixed Slopes (Between)",
                                "Fixed Slopes (Within)")),
                 y = as.vector(mod1a.r2$result$rs$decomp))

# Draw bar chart in line with the default setting of multilevel.r2()
ggplot(df, aes(x = var, y = y, fill = part)) +

```

```

theme_bw() +
geom_bar(stat = "identity") +
scale_fill_manual(values = c("#E69F00", "#009E73", "#CC79A7", "#0072B2", "#D55E00")) +
scale_y_continuous(name = "Proportion of Variance", breaks = seq(0, 1, by = 0.1)) +
theme(axis.title.x = element_blank(),
      axis.ticks.x = element_blank(),
      legend.title = element_blank(),
      legend.position = "bottom",
      legend.box.margin = margin(-10, 6, 6, 6)) +
guides(fill = guide_legend(nrow = 2, reverse = TRUE))

## End(Not run)

```

multilevel.r2.manual	<i>R-Squared Measures for Multilevel and Linear Mixed Effects Models by Rights and Sterba (2019), Manually Inputting Parameter Estimates</i>
----------------------	--

Description

This function computes R-squared measures by Rights and Sterba (2019) for multilevel and linear mixed effects models by manually inputting parameter estimates.

Usage

```

multilevel.r2.manual(data, within = NULL, between = NULL, random = NULL,
  gamma.w = NULL, gamma.b = NULL, tau, sigma2,
  intercept = TRUE, center = TRUE, digits = 3,
  plot = FALSE, gray = FALSE, start = 0.15, end = 0.85,
  color = c("#D55E00", "#0072B2", "#CC79A7", "#009E73", "#E69F00"),
  filename = NULL, width = NA, height = NA,
  units = c("in", "cm", "mm", "px"), dpi = 600,
  write = NULL, append = TRUE, check = TRUE, output = TRUE)

```

Arguments

data	a matrix or data frame with the level-1 and level-2 predictors and outcome variable used in the model.
within	a character vector with the variable names in data or numeric vector with numbers corresponding to the columns in data of the level-1 predictors used in the model. If none used, set to NULL.
between	a character vector with the variable names in data or numeric vector with numbers corresponding to the columns in data of the level-2 predictors used in the model. If none used, set to NULL.
random	a character vector with the variable names in data or numeric vector with numbers corresponding to the columns in data of the level-1 predictors that have random slopes in the model. If no random slopes specified, set to NULL.

<code>gamma.w</code>	a numeric vector of fixed slope estimates for all level-1 predictors, to be entered in the order of the predictors listed in the argument <code>within</code> .
<code>gamma.b</code>	a numeric vector of the intercept and fixed slope estimates for all level-2 predictors, to be entered in the order of the predictors listed in the argument <code>between</code> . Note that the first element is the parameter estimate for the intercept if <code>intercept = TRUE</code> .
<code>tau</code>	a matrix indicating the random effects covariance matrix, the first row/column denotes the intercept variance and covariances (if <code>intercept</code> is fixed, set all to 0) and each subsequent row/column denotes a given random slope's variance and covariances (to be entered in the order listed in the argument <code>random</code>).
<code>sigma2</code>	a numeric value indicating the level-1 residual variance.
<code>intercept</code>	logical: if <code>TRUE</code> (default), the first element in the <code>gamma.b</code> is assumed to be the fixed intercept estimate; if set to <code>FALSE</code> , the first element in the argument <code>gamma.b</code> is assumed to be the first fixed level-2 predictor slope.
<code>center</code>	logical: if <code>TRUE</code> (default), all level-1 predictors are assumed to be cluster-mean-centered and the function will output all decompositions; if set to <code>FALSE</code> , function will output only the total decomposition.
<code>digits</code>	an integer value indicating the number of decimal places to be used.
<code>plot</code>	logical: if <code>TRUE</code> , bar chart showing the decomposition of scaled total, within-cluster, and between-cluster outcome variance into five (total), three (within-cluster), and two (between-cluster) proportions is drawn. Note that the ggplot2 package is required to draw the bar chart.
<code>gray</code>	logical: if <code>TRUE</code> , graphical parameter to draw the bar chart in gray scale.
<code>start</code>	a numeric value between 0 and 1, graphical parameter to specify the gray value at the low end of the palette.
<code>end</code>	a numeric value between 0 and 1, graphical parameter to specify the gray value at the high end of the palette.
<code>color</code>	a character vector, graphical parameter indicating the color of bars in the bar chart in the following order: Fixed slopes (Within), Fixed slopes (Between), Slope variation (Within), Intercept variation (Between), and Residual (Within). By default, colors from the colorblind-friendly palettes are used.
<code>filename</code>	a character string indicating the filename argument including the file extension in the <code>ggsave</code> function. Note that one of <code>".eps"</code> , <code>".ps"</code> , <code>".tex"</code> , <code>".pdf"</code> (default), <code>".jpeg"</code> , <code>".tiff"</code> , <code>".png"</code> , <code>".bmp"</code> , <code>".svg"</code> or <code>".wmf"</code> needs to be specified as file extension in the file argument. Note that plots can only be saved when <code>plot = TRUE</code> .
<code>width</code>	a numeric value indicating the width argument (default is the size of the current graphics device) in the <code>ggsave</code> function.
<code>height</code>	a numeric value indicating the height argument (default is the size of the current graphics device) in the <code>ggsave</code> function.
<code>units</code>	a character string indicating the units argument (default is <code>in</code>) in the <code>ggsave</code> function.
<code>dpi</code>	a numeric value indicating the dpi argument (default is 600) in the <code>ggsave</code> function.

<code>write</code>	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

A number of R-squared measures for multilevel and linear mixed effects models have been developed in the methodological literature (see Rights & Sterba, 2018). R-squared measures by Rights and Sterba (2019) provide an integrative framework of R-squared measures for multilevel and linear mixed effects models with random intercepts and/or slopes. Their measures are based on partitioning model implied variance from a single fitted model, but they provide a full partitioning of the total outcome variance to one of five specific sources. See the help page of the [multilevel.r2](#) function for more details.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	matrix or data frame specified in <code>data</code>
<code>plot</code>	ggplot2 object for plotting the results
<code>args</code>	specification of function arguments
<code>result</code>	list with result tables, i.e., <code>decomp</code> for the decomposition, <code>total</code> for total R2 measures, <code>within</code> for the within-cluster R2 measures, and <code>between</code>

for the between-cluster R2 measures.

Note

This function is based on a copy of the function `r2mlm_manual()` in the **r2mlm** package by Mairead Shaw, Jason Rights, Sonya Sterba, and Jessica Flake.

Author(s)

Jason D. Rights, Sonya K. Sterba, Jessica K. Flake, and Takuya Yanagida

References

- Rights, J. D., & Cole, D. A. (2018). Effect size measures for multilevel models in clinical child and adolescent research: New r-squared methods and recommendations. *Journal of Clinical Child and Adolescent Psychology*, 47, 863-873. <https://doi.org/10.1080/15374416.2018.1528550>
- Rights, J. D., & Sterba, S. K. (2019). Quantifying explained variance in multilevel models: An integrative framework for defining R-squared measures. *Psychological Methods*, 24, 309-338. <https://doi.org/10.1037/met0000184>

See Also

[multilevel.r2](#), [multilevel.cor](#), [multilevel.descript](#), [multilevel.icc](#), [multilevel.indirect](#)

Examples

```
## Not run:

# Load misty and lme4 package
misty::libraries(misty, lme4)

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

#-----

Demo.twolevel <- center(Demo.twolevel, x2, type = "CWC", cluster = "cluster")

# Compute group means, cluster.scores() from the misty package
Demo.twolevel <- cluster.scores(Demo.twolevel, x2, cluster = "cluster", name = "x2.b")

# Estimate random intercept model using the lme4 package
mod1 <- lmer(y1 ~ x2.c + x2.b + w1 + (1| cluster), data = Demo.twolevel,
            REML = FALSE, control = lmerControl(optimizer = "bobyqa"))

# Estimate random intercept and slope model using the lme4 package
mod2 <- lmer(y1 ~ x2.c + x2.b + w1 + (1 + x2.c | cluster), data = Demo.twolevel,
            REML = FALSE, control = lmerControl(optimizer = "bobyqa"))

#-----

# Example 1: Random intercept model

# Fixed slope estimates
fixef(mod1)

# Random effects variance-covariance matrix
as.data.frame(VarCorr(mod1))

# R-squared measures according to Rights and Sterba (2019)
multilevel.r2.manual(data = Demo.twolevel,
                    within = "x2.c", between = c("x2.b", "w1"),
                    gamma.w = 0.41127956,
                    gamma.b = c(0.01123245, -0.08269374, 0.17688507),
                    tau = 0.9297401,
                    sigma2 = 1.813245794)

#-----

# Example 2: Random intercept and slope model

# Fixed slope estimates
fixef(mod2)

# Random effects variance-covariance matrix
```

```

as.data.frame(VarCorr(mod2))

# R-squared measures according to Rights and Sterba (2019)
multilevel.r2.manual(data = Demo.twolevel,
  within = "x2.c", between = c("x2.b", "w1"), random = "x2.c",
  gamma.w = 0.41127956,
  gamma.b = c(0.01123245, -0.08269374, 0.17688507),
  tau = matrix(c(0.931008649, 0.004110479, 0.004110479, 0.017068857), ncol = 2),
  sigma2 = 1.813245794)

## End(Not run)

```

na.as	<i>Replace Missing Values With User-Specified Values or User-Specified Values With Missing Values</i>
-------	---

Description

The function `na.as` replaces NA in a vector, factor, list, matrix or data frame with a user-specified value or character string in the argument `na`, while the function `as.na` replaces user-specified values in the argument `na` in a vector, factor, matrix, array, list, or data frame with NA.

Usage

```

na.as(data, ..., na, replace = TRUE, as.na = NULL, check = TRUE)

as.na(data, ..., na, replace = TRUE, check = TRUE)

```

Arguments

<code>data</code>	a vector, factor, matrix, array, data frame, or list.
<code>...</code>	an expression indicating the variable names in <code>data</code> , e.g., <code>as.na(dat, x1, x2)</code> for selecting the variables <code>x1</code> and <code>x2</code> from the data frame <code>dat</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>na</code>	a vector indicating values or characters to replace with NA, or which NA is replaced. Note that a numeric value or character string needs to be specified for the argument <code>na</code> when using <code>na.as</code> .
<code>replace</code>	logical: if TRUE (default), variable(s) specified in <code>...</code> are replaced in the argument <code>data</code> .
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>as.na</code>	a numeric vector or character vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.

Value

Returns a vector, factor, matrix, array, data frame, or list specified in the argument `data`.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Becker, R. A., Chambers, J. M. and Wilks, A. R. (1988) *The New S Language*. Wadsworth & Brooks/Cole.

See Also

[na.auxiliary](#), [na.coverage](#), [na.descript](#), [na.indicator](#), [na.pattern](#), [na.prop](#), [na.test](#)

Examples

```
#-----
# Numeric vector
num <- c(1, 3, 2, 4, 5)

# Example 1i: Replace NA with 2
na.as(c(1, 3, NA, 4, 5), na = 2)

# Example 1b: Replace 2 with NA
as.na(num, na = 2)

# Example 1c: Replace 2, 3, and 4 with NA
as.na(num, na = c(2, 3, 4))

#-----
# Character vector
chr <- c("a", "b", "c", "d", "e")

# Example 2a: Replace NA with "b"
na.as(c("a", NA, "c", "d", "e"), na = "b")

# Example 2b: Replace "b" with NA
as.na(chr, na = "b")

# Example 2c: Replace "b", "c", and "d" with NA
as.na(chr, na = c("b", "c", "d"))

#-----
# Factor
fac <- factor(c("a", "a", "b", "b", "c", "c"))

# Example 3a: Replace NA with "b"
na.as(factor(c("a", "a", NA, NA, "c", "c")), na = "b")

# Example 3b: Replace "b" with NA
as.na(fac, na = "b")

# Example 3c: Replace "b" and "c" with NA
as.na(fac, na = c("b", "c"))
```

```

#-----
# Matrix
mat <- matrix(1:20, ncol = 4)

# Example 4a: Replace NA with 2
na.as(matrix(c(1, NA, 3, 4, 5, 6), ncol = 2), na = 2)

# Example 4b: Replace 8 with NA
as.na(mat, na = 8)

# Example 4c: Replace 8, 14, and 20 with NA
as.na(mat, na = c(8, 14, 20))

#-----
# Array

# Example 5: Replace 1 and 10 with NA
as.na(array(1:20, dim = c(2, 3, 2)), na = c(1, 10))

#-----
# List

# Example 6: Replace 1 with NA
as.na(list(x1 = c(1, 2, 3, 1, 2, 3), x2 = c(2, 1, 3, 2, 1)), na = 1)

#-----
# Data frame
df <- data.frame(x1 = c(1, NA, 3), x2 = c(2, 1, 3), x3 = c(3, NA, 2))

# Example 7a: Replace NA with -99
na.as(df, na = -99)

# Example 7b: Replace 1 with NA
as.na(df, na = 1)

# Example 7c: Replace 1 with NA for the variable 'x2'
as.na(df, x2, na = 1)

# Alternative specification
as.na(df$x2, na = 1)

# Example 7d: Replace 1 and 3 with NA
as.na(df, na = c(1, 3))

# Example 7e: Replace 1 with NA in 'x2' and 'x3'
as.na(df, x2, x3, na = 1)

```

Description

This function computes (1) a matrix with Pearson product-moment correlation for continuous variables, multiple correlation coefficient for categorical and continuous variables, and Phi coefficient and Cramer's V for categorical variables to identify variables related to the incomplete variable (i.e., correlates of incomplete variables), (2) a matrix with Cohen's d , Phi coefficient and Cramer's V for comparing cases with and without missing values, and (3) semi-partial correlations of an outcome variable conditional on the predictor variables of a substantive model with a set of candidate auxiliary variables to identify correlates of an incomplete outcome variable as suggested by Raykov and West (2016).

Usage

```
na.auxiliary(data, ..., model = NULL, categ = NULL, estimator = c("ML", "MLR"),
             missing = c("fiml", "two.stage", "robust.two.stage", "doubly.robust"),
             adjust = TRUE, weighted = FALSE, correct = FALSE,
             tri = c("both", "lower", "upper"), digits = 2, p.digits = 3,
             as.na = NULL, write = NULL, append = TRUE,
             check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a data frame with incomplete data, where missing values are coded as NA.
<code>...</code>	an expression indicating the variable names in <code>data</code> , e.g., <code>na.auxiliary(dat, x1, x2, x3)</code> . Categorical variables specified in the argument <code>categ</code> can be, but do not need to be selected using the <code>...</code> argument. Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>model</code>	a character string specifying the substantive model predicting a continuous outcome variable using a set of predictor variables to estimate semi-partial correlations between the outcome variable and a set of candidate auxiliary variables. The default setting is <code>model = NULL</code> , i.e., the function computes Pearson product-moment correlation matrix and Cohen's d matrix.
<code>categ</code>	a character vector specifying the variables that are treated as categorical (see 'Details'). Note that variables that are factors or character vectors will be automatically added to the argument <code>categ</code> . Categorical variables will be excluded from the analysis when specifying the <code>model</code> argument to compute semi-partial correlations.
<code>estimator</code>	a character string indicating the estimator to be used when estimating semi-partial correlation coefficients, i.e., "ML" for maximum likelihood parameter estimates with conventional standard errors or "MLR" (default) maximum likelihood parameter estimates with Huber-White robust standard errors.
<code>missing</code>	a character string indicating how to deal with missing data when estimating semi-partial correlation coefficients, i.e., "fiml" for full information maximum likelihood method, <code>two.stage</code> for two-stage maximum likelihood method, <code>robust.two.stage</code> for robust two-stage maximum likelihood method, and <code>doubly-robust</code> for doubly-robust method (see 'Details' in the item.cfa function). The default setting is <code>missing = "fiml"</code> .

adjust	logical: if TRUE (default), phi coefficient is adjusted by relating the coefficient to the possible maximum and Cramer's V is corrected for small-sample bias.
weighted	logical: if TRUE (default), the weighted pooled standard deviation is used when computing Cohen's d .
correct	logical: if TRUE, correction factor for Cohen's d to remove positive bias in small samples is used.
tri	a character string indicating which triangular of the correlation matrix to show on the console, i.e., both for upper and lower triangular, lower (default) for the lower triangular, and upper for the upper triangular.
digits	integer value indicating the number of decimal places digits to be used for displaying correlation coefficients and Cohen's d estimates.
p.digits	an integer value indicating the number of decimal places to be used for displaying the p -value.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

The function computes matrices with statistical measures depending on the level of measurement of the variables involved in the analysis:

Variables Related to the Incomplete Variable *Continuous variables*: Product-moment correlation coefficient is computed for continuous variables.

- *Continuous and categorical variable*: Multiple correlation coefficient (R) is computed based on a linear model with a dummy-coded categorical variable as predictor, where the multiple correlation coefficient is the square root of the coefficient of determination of this model. Note that the multiple R for a binary predictor variable is equivalent to the point-biserial correlation coefficient between the binary variable and the continuous outcome.
- *Categorical variables*: Phi coefficient is computed for two dichotomous variables, while Cramer's V is computed when one of the categorical variables is polytomous

Variables Related to the Probability of Missigness • *Continuous variable*: Cohen's d is computed to investigate mean differences in the continuous variable depending on cases with and without missing values.

- *Categorical variable*: Phi coefficient is computed to investigate the association between the grouping variable (0 = observed, 1 = missing) and a dichotomous variable, while Cramer's V is computed when the categorical variable is polytomous.

Substantive model predicting a continuous outcome variable Categorical variables are removed before computing semi-partial correlations based on the approach suggested by Raykov and West (2016).

Note that factors and characters are treated as categorical variables regardless of the specification of the argument `categ`, while numeric vectors in the data frame are treated as continuous variables if they are not specified in the argument `categ`.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame used for the current analysis
<code>model</code>	lavaan model syntax for estimating the semi-partial correlations
<code>model.fit</code>	fitted lavaan model for estimating the semi-partial correlations
<code>args</code>	pecification of function arguments
<code>result</code>	list with result tables

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Enders, C. K. (2022). *Applied missing data analysis* (2nd ed.). The Guilford Press.
- Graham, J. W. (2009). Missing data analysis: Making it work in the real world. *Annual Review of Psychology*, 60, 549-576. <https://doi.org/10.1146/annurev.psych.58.110405.085530>
- Raykov, T., & West, B. T. (2016). On enhancing plausibility of the missing at random assumption in incomplete data analyses via evaluation of response-auxiliary variable correlations. *Structural Equation Modeling*, 23(1), 45–53. <https://doi.org/10.1080/10705511.2014.937848>
- van Buuren, S. (2018). *Flexible imputation of missing data* (2nd ed.). Chapman & Hall.

See Also

[as.na](#), [na.as](#), [na.coverage](#), [na.descript](#), [na.indicator](#), [na.pattern](#), [na.prop](#), [na.test](#)

Examples

```
# Example 1a: Auxiliary variables
na.auxiliary(airquality)

# Example 1b: Auxiliary variables, "Month" as categorical variable
na.auxiliary(airquality, categ = "Month")

# Example 2: Semi-partial correlation coefficients
na.auxiliary(airquality, model = "Ozone ~ Solar.R + Wind")
```

```
## Not run:
# Example 3a: Write Results into a text file
na.auxiliary(airquality, write = "NA_Auxiliary.txt")

# Example 3a: Write Results into an Excel file
na.auxiliary(airquality, write = "NA_Auxiliary.xlsx")

## End(Not run)
```

na.coverage

Variance-Covariance Coverage

Description

This function computes the proportion of cases that contributes for the calculation of each variance and covariance.

Usage

```
na.coverage(data, ..., tri = c("both", "lower", "upper"), digits = 2,
            as.na = NULL, write = NULL, append = TRUE, check = TRUE,
            output = TRUE)
```

Arguments

<code>data</code>	a data frame with incomplete data, where missing values are coded as NA.
<code>...</code>	an expression indicating the variable names in data, e.g., <code>na.coverage(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>tri</code>	a character string or character vector indicating which triangular of the matrix to show on the console, i.e., both for upper and lower triangular, lower (default) for the lower triangular, and upper for the upper triangular.
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying proportions.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension <code>".txt"</code> (e.g., <code>"Output.txt"</code>) or Excel file with file extension <code>".xlsx"</code> (e.g., <code>"Output.xlsx"</code>). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension <code>.txt</code> specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame used for the current analysis
<code>args</code>	specification of function arguments
<code>result</code>	result table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Enders, C. K. (2022). *Applied missing data analysis* (2nd ed.). The Guilford Press.
- Graham, J. W. (2009). Missing data analysis: Making it work in the real world. *Annual Review of Psychology*, 60, 549-576. <https://doi.org/10.1146/annurev.psych.58.110405.085530>
- van Buuren, S. (2018). *Flexible imputation of missing data* (2nd ed.). Chapman & Hall.

See Also

[write.result](#), [as.na](#), [na.as](#), [na.auxiliary](#), [na.descript](#), [na.indicator](#), [na.pattern](#), [na.prop](#), [na.test](#)

Examples

```
# Example 1: Compute variance-covariance coverage
na.coverage(airquality)

## Not run:
# Example 2a: Write Results into a text file
na.coverage(airquality, write = "Coverage.txt")

# Example 2b: Write Results into a Excel file
na.coverage(airquality, write = "Coverage.xlsx")

## End(Not run)
```

na.descript	<i>Descriptive Statistics for Missing Data in Single-Level, Two-Level and Three-Level Data</i>
-------------	--

Description

This function computes descriptive statistics for missing data in single-level, two-level, and three-level data, e.g. number of incomplete cases, number of missing values, and summary statistics for the number of missing values across all variables.

Usage

```
na.descript(data, ..., cluster = NULL, table = FALSE, digits = 2,
            as.na = NULL, write = NULL, append = TRUE, check = TRUE,
            output = TRUE)
```

Arguments

data	a data frame with incomplete data, where missing values are coded as NA.
...	an expression indicating the variable names in data, e.g., <code>na.descript(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the <code>df.subset</code> function.
cluster	a character string indicating the name of the cluster variable in data for two-level data, a character vector indicating the names of the cluster variables in data for three-level data, or a vector or data frame representing the nested grouping structure (i.e., group or cluster variables). Alternatively, a character string or character vector indicating the variable name(s) of the cluster variable(s) in data. Note that the cluster variable at Level 3 come first in a three-level model, i.e., <code>cluster = c("level3", "level2")</code> .
table	logical: if TRUE, a frequency table with number of observed values (" <code>nOb</code> "), percent of observed values (" <code>pOb</code> "), number of missing values (" <code>nNA</code> "), and percent of missing values (" <code>pNA</code> ") is printed for each variable on the console.
digits	an integer value indicating the number of decimal places to be used for displaying percentages.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
write	a character string naming a file for writing the output into either a text file with file extension " <code>.txt</code> " (e.g., " <code>Output.txt</code> ") or Excel file with file extension " <code>.xlsx</code> " (e.g., " <code>Output.xlsx</code> "). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension <code>.txt</code> specified in <code>write</code> , if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	data frame used for the current analysis
<code>args</code>	specification of function arguments
<code>result</code>	list with results

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Enders, C. K. (2022). *Applied missing data analysis* (2nd ed.). The Guilford Press.
- Graham, J. W. (2009). Missing data analysis: Making it work in the real world. *Annual Review of Psychology*, 60, 549-576. <https://doi.org/10.1146/annurev.psych.58.110405.085530>
- van Buuren, S. (2018). *Flexible imputation of missing data* (2nd ed.). Chapman & Hall.

See Also

[write.result](#), [as.na](#), [na.as](#), [na.auxiliary](#), [na.coverage](#), [na.indicator](#), [na.pattern](#), [na.prop](#), [na.test](#)

Examples

```
#-----
# Single-Level Data

# Example 1: Descriptive statistics for missing data
na.descript(airquality)

# Example 2: Descriptive statistics for missing data, print results with 3 digits
na.descript(airquality, digits = 3)

# Example 3: Descriptive statistics for missing data with frequency table
na.descript(airquality, table = TRUE)

#-----
# Two-Level Data

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

# Example 4: Descriptive statistics for missing data
na.descript(Demo.twolevel, cluster = "cluster")

#-----
```

```
# Three-Level Data

# Create arbitrary three-level data
Demo.threelevel <- data.frame(Demo.twolevel, cluster2 = Demo.twolevel$cluster,
                             cluster3 = rep(1:10, each = 250))

# Example 5: Descriptive statistics for missing data
na.descript(Demo.threelevel, cluster = c("cluster3", "cluster2"))

## Not run:
#-----
# Write Results

# Example 6a: Write Results into a text file
na.descript(airquality, table = TRUE, write = "NA_Descriptives.txt")

# Example 6b: Write Results into a Excel file
na.descript(airquality, table = TRUE, write = "NA_Descriptives.xlsx")

## End(Not run)
```

na.indicator	<i>Missing Data Indicator Matrix</i>
--------------	--------------------------------------

Description

This function creates a missing data indicator matrix R that denotes whether values are observed or missing, i.e., $r = 0$ if a value is observed, and $r = 1$ if a value is missing.

Usage

```
na.indicator(data, ..., na = 1, append = TRUE, name = ".i", as.na = NULL,
             check = TRUE)
```

Arguments

data	a data frame with incomplete data, where missing values are coded as NA.
...	an expression indicating the variable names in data, e.g., <code>na.indicator(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
na	an integer value specifying the value representing missing values, i.e., either <code>na = 0</code> for <code>0 = missing</code> and <code>1 = observed</code> , or <code>na = 1</code> (default) for <code>0 (observed)</code> and <code>1 = missing</code> .
append	logical: if TRUE (default), missing data indicator matrix is appended to the data frame specified in the argument data.
name	a character string indicating the name suffix of indicator variables By default, the indicator variables are named with the ending <code>".i"</code> resulting in e.g. <code>"x1.i"</code> and <code>"x2.i"</code> . Note that when selecting one single variable, the indicator variable is named <code>x.i</code> by default or named after the argument name.

- as.na a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
- check logical: if TRUE (default), argument specification is checked.

Value

Returns a matrix or data frame with $r = 1$ if a value is observed, and $r = 0$ if a value is missing.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Enders, C. K. (2022). *Applied missing data analysis* (2nd ed.). The Guilford Press.

Graham, J. W. (2009). Missing data analysis: Making it work in the real world. *Annual Review of Psychology*, 60, 549-576. <https://doi.org/10.1146/annurev.psych.58.110405.085530>

van Buuren, S. (2018). *Flexible imputation of missing data* (2nd ed.). Chapman & Hall.

See Also

[as.na](#), [na.as](#), [na.auxiliary](#), [na.coverage](#), [na.descript](#), [na.pattern](#), [na.prop](#), [na.test](#)

Examples

```
# Example 1: Create missing data indicator matrix
na.indicator(airquality)

# Example 2: Do not append missing data indicator matrix to the data frame
na.indicator(airquality, append = FALSE)
```

na.pattern	<i>Missing Data Pattern</i>
------------	-----------------------------

Description

This function computes a summary of missing data patterns, i.e., number (cases with a specific missing data pattern and plots the missing data patterns.

Usage

```
na.pattern(data, ..., order = FALSE, n.pattern = NULL, digits = 2, as.na = NULL,
  plot = FALSE, square = TRUE, rotate = FALSE,
  color = c("#B61A51B3", "#006CC2B3"), tile.alpha = 0.6,
  plot.margin = c(4, 16, 0, 4), legend.box.margin = c(-8, 6, 6, 6),
  legend.key.size = 12, legend.text.size = 9, filename = NULL,
  width = NA, height = NA, units = c("in", "cm", "mm", "px"),
  dpi = 600, write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

<code>data</code>	a data frame with incomplete data, where missing values are coded as NA.
<code>...</code>	an expression indicating the variable names in data e.g., <code>na.pattern(dat, x1, x2, x3)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>order</code>	logical: if TRUE, variables are ordered from left to right in increasing order of missing values.
<code>n.pattern</code>	an integer value indicating the minimum number of cases sharing a missing data pattern to be included in the result table and the plot, e.g., specifying <code>n.pattern = 5</code> excludes missing data patterns with less than 5 cases.
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying percentages.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>plot</code>	logical: if TRUE, missing data pattern is plotted.
<code>square</code>	logical: if TRUE (default), the plot tiles are squares to mimic the <code>md.pattern</code> function in the package mice .
<code>rotate</code>	logical: if TRUE, the variable name labels are rotated 90 degrees.
<code>color</code>	a character string indicating the color for the "fill" argument. Note that the first color represents missing values and the second color represent observed values.
<code>tile.alpha</code>	a numeric value between 0 and 1 for the alpha argument (default is 0.1).
<code>plot.margin</code>	a numeric vector indicating the <code>plot.margin</code> argument for the theme function.
<code>legend.box.margin</code>	a numeric vector indicating the <code>legend.box.margin</code> argument for the theme function.
<code>legend.key.size</code>	a numeric value indicating the <code>legend.key</code> argument (default is <code>unit(12, "pt")</code>) for the theme function.
<code>legend.text.size</code>	a numeric value indicating the <code>legend.text</code> argument (default is <code>element_text(size = 10)</code>) for the theme function.
<code>filename</code>	a character string indicating the <code>filename</code> argument (default is "NA_Pattern.pdf") including the file extension for the <code>ggsave</code> function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the <code>file</code> argument.
<code>width</code>	a numeric value indicating the <code>width</code> argument (default is the size of the current graphics device) for the <code>ggsave</code> function.
<code>height</code>	a numeric value indicating the <code>height</code> argument (default is the size of the current graphics device) for the <code>ggsave</code> function.
<code>units</code>	a character string indicating the <code>units</code> argument (default is <code>in</code>) for the <code>ggsave</code> function.

dpi	a numeric value indicating the dpi argument (default is 600) for the ggsave function.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	data frame with variables used in the analysis
args	specification of function arguments
result	result table
plot	ggplot2 object for plotting the results
pattern	a numeric vector indicating the missing data pattern for each case

Note

The code for plotting missing data patterns is based on the `plot_pattern` function in the **ggmice** package by Hanne Oberman.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Enders, C. K. (2022). *Applied missing data analysis* (2nd ed.). The Guilford Press.
- Graham, J. W. (2009). Missing data analysis: Making it work in the real world. *Annual Review of Psychology*, 60, 549-576. <https://doi.org/10.1146/annurev.psych.58.110405.085530>
- Oberman, H. (2023). *ggmice: Visualizations for 'mice' with 'ggplot2'*. R package version 0.1.0. <https://doi.org/10.32614/CRAN.package.ggmice>
- van Buuren, S. (2018). *Flexible imputation of missing data* (2nd ed.). Chapman & Hall.

See Also

[write.result](#), [as.na](#), [na.as](#), [na.auxiliary](#), [na.coverage](#), [na.descript](#), [na.indicator](#), [na.prop](#), [na.test](#)

Examples

```
# Example 1: Compute a summary of missing data patterns
dat.pattern <- na.pattern(airquality)

# Example 2a: Compute and plot a summary of missing data patterns
na.pattern(airquality, plot = TRUE)

# Example 2b: Exclude missing data patterns with less than 3 cases
na.pattern(airquality, plot = TRUE, n.pattern = 3)

# Example 3: Vector of missing data pattern for each case
dat.pattern$pattern

# Data frame without cases with missing data pattern 2 and 4
airquality[!dat.pattern$pattern == 2, ]

## Not run:
# Example 4a: Write Results into a text file
na.pattern(airquality, write = "NA_Pattern.xlsx")

# Example 4b: Write Results into a Excel file
na.pattern(airquality, write = "NA_Pattern.xls")

## End(Not run)
```

na.prop	<i>Proportion of Missing Data for Each Case</i>
---------	---

Description

This function computes the proportion of missing data for each case in a data frame.

Usage

```
na.prop(data, ..., digits = 2, append = TRUE, name = "na.prop",
        as.na = NULL, check = TRUE)
```

Arguments

data	a data frame with incomplete data, where missing values are coded as NA.
...	an expression indicating the variable names in data, e.g., na.prop(dat, x1, x2, x3). Note that the operators +, -, ~, :, ::, and ! can also be used to select variables, see 'Details' in the df.subset function.
name	a character string indicating the name of the variable appended to the data frame specified in the argument data when append = TRUE.
.	

append	logical: if TRUE (default), variable with proportion of missing data is appended to the data frame specified in the argument data
digits	an integer value indicating the number of decimal places to be used for displaying proportions.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
check	logical: if TRUE, argument specification is checked.

Value

Returns a numeric vector with the same length as the number of rows in data containing the proportion of missing data.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Enders, C. K. (2022). *Applied missing data analysis* (2nd ed.). The Guilford Press.
- Graham, J. W. (2009). Missing data analysis: Making it work in the real world. *Annual Review of Psychology*, 60, 549-576. <https://doi.org/10.1146/annurev.psych.58.110405.085530>
- van Buuren, S. (2018). *Flexible imputation of missing data* (2nd ed.). Chapman & Hall.

See Also

[as.na](#), [na.as](#), [na.auxiliary](#), [na.coverage](#), [na.descript](#), [na.indicator](#), [na.pattern](#), [na.test](#)

Examples

```
# Example 1: Compute proportion of missing data for each case in the data frame
na.prop(airquality)

# Example 2: Do not append proportions of missing data to the data frame
na.prop(airquality, append = FALSE)
```

na.satcor

Fit a Saturated Correlates Model

Description

This function estimates a confirmatory factor analysis model (`cfa.satcor` function), structural equation model (`sem.satcor` function), growth curve model (`growth.satcor` function), or latent variable model (`lavaan.satcor` function) in the R package **lavaan** using full information maximum likelihood (FIML) method to handle missing data while automatically specifying a saturated correlates model to incorporate auxiliary variables into a substantive model without affecting the parameter estimates, the standard errors, or the estimates of quality of fit (Graham, 2003).

Usage

```
na.satcor(model, data, aux, fun = c("cfa", "sem", "growth", "lavaan"),
          check = TRUE, ...)
```

```
cfa.satcor(model, data, aux, check = TRUE, ...)
```

```
sem.satcor(model, data, aux, check = TRUE, ...)
```

```
growth.satcor(model, data, aux, check = TRUE, ...)
```

```
lavaan.satcor(model, data, aux, check = TRUE, ...)
```

Arguments

model	a character string indicating the lavaan model syntax without the auxiliary variables specified in aux.
data	a data frame containing the observed variables used in the lavaan model syntax specified in model and the auxiliary variables specified in aux.
aux	a character vector indicating the names of the auxiliary variables in the data frame specified in data that will be added to the lavaan model syntax specified in model. Note that this function can only incorporate continuous auxiliary variables, i.e., the function cannot deal with categorical auxiliary variables.
fun	a character string indicating the name of a specific lavaan function used to fit model, i.e., cfa, sem, growth, or lavaan. Note that this argument is only required for the function na.satcor.
check	logical: if TRUE (default), argument specification is checked.
...	additional arguments passed to the lavaan function.

Value

An object of class lavaan, for which several methods are available in the R package **lavaan**, including a summary method.

Note

This function is a modified copy of the auxiliary(), cfa.auxiliary(), sem.auxiliary(), growth.auxiliary(), and lavaan.auxiliary() functions in the **semTools** package by Terrence D. Jorgensen et al. (2022).

Author(s)

Takuya Yanagida

References

Graham, J. W. (2003). Adding missing-data-relevant variables to FIML-based structural equation models. *Structural Equation Modeling*, 10(1), 80-100. https://doi.org/10.1207/S15328007SEM1001_4

Jorgensen, T. D., Pornprasertmanit, S., Schoemann, A. M., & Rosseel, Y. (2022). *semTools: Useful tools for structural equation modeling*. R package version 0.5-6. Retrieved from <https://CRAN.R-project.org/package=semTools>

Examples

```
# Load lavaan package
library(lavaan)

#-----
# Example 1: Saturated correlates model for the sem function

# Model specification
model <- 'Ozone ~ Wind'

# Model estimation using the sem.satcor function
mod.fit <- sem.satcor(model, data = airquality, aux = c("Temp", "Month"))

# Model estimation using the na.satcor function
mod.fit <- na.satcor(model, data = airquality, fun = "sem", aux = c("Temp", "Month"),
  estimator = "MLR")

# Result summary
summary(mod.fit)
```

na.test

Missing Completely at Random (MCAR) Test

Description

This function performs Little's Missing Completely at Random (MCAR) test and Jamshidian and Jalal's approach for testing the MCAR assumption. By default, the function performs the Little's MCAR test.

Usage

```
na.test(data, ..., print = c("all", "little", "jamjal"),
  impdat = NULL, delete = 6, method = c("npar", "normal"),
  m = 20, seed = 123, nrep = 10000, n.min = 30,
  pool = c("m", "med", "min", "max", "random"),
  alpha = 0.05, digits = 2, p.digits = 3, as.na = NULL,
  write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

data	a data frame with incomplete data, where missing values are coded as NA.
...	an expression indicating the variable names in data, e.g., <code>na.test(dat, x1, x2, x3)</code> . Note that the operators <code>.</code> , <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the <code>df.subset</code> function.

print	a character vector indicating which results to be printed on the console, i.e. "all" for Little's MCAR test and Jamshidian and Jalal's approach, "little" (default) for Little's MCAR test, and "jamjal" for Jamshidian and Jalal's approach.
impdat	an object of class <code>mids</code> from the mice package to provide a data set multiply imputed in the mice package. The function will not impute the data data set specified in the argument <code>data</code> when specifying this argument and will use the imputed data sets provided in the argument <code>impdat</code> for performing the Jamshidian and Jalal's approach. Note that the argument <code>data</code> still needs to be specified because the variables used for the analysis are extracted from the data frame specified in <code>data</code> .
delete	an integer value indicating missing data patterns consisting of delete number of cases or less removed from the Jamshidian and Jalal's approach. The default setting is <code>delete = 6</code> .
method	a character string indicating the imputation method, i.e., "npar" for using a non-parametric imputation method by Sirvastava and Dolatabadi (2009) or "normal" for imputing missing data assuming that the data come from a multivariate normal distribution (see Jamshidian & Jalal, 2010).
m	an integer value indicating the number of multiple imputations. The default setting is <code>m = 20</code> .
seed	an integer value that is used as argument by the <code>set.seed</code> function for off-setting the random number generator before performing Jamshidian and Jalal's approach. The default setting is <code>seed = 123</code> . Set the value to <code>NULL</code> to specify a system selected seed.
nrep	an integer value indicating the replications used to simulate the Neyman distribution to determine the cut off value for the Neyman test. Larger values increase the accuracy of the Neyman test. The default setting is <code>nrep = 10000</code> .
n.min	an integer value indicating the minimum number of cases in a group that triggers the use of asymptotic Chi-square distribution in place of the empirical distribution in the Neyman test of uniformity.
pool	a character string indicating the pooling method, i.e., "m" for computing the average test statistic and p-values, "med" for computing the median test statistic and p-values, "min" for computing the maximum test statistic and minimum p-values, "max" for computing the minimum test statistic and maximum p-values, and "random" for randomly choosing a test statistic and corresponding p-value from repeated complete data analyses. The default setting is <code>pool = "med"</code> , i.e., median test statistic and p-values are computed as suggested by Eekhout, Wiel and Heymans (2017).
alpha	a numeric value between 0 and 1 indicating the significance level of the Hawkins test. The default setting is <code>alpha = 0.05</code> , i.e., the Anderson-Darling non-parametric test is provided when the p-value of the Hawkins test is less than or equal <code>0.05</code> .
digits	an integer value indicating the number of decimal places to be used for displaying results.
p.digits	an integer value indicating the number of decimal places to be used for displaying the <i>p</i> -value.

as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown.

Details

Little's MCAR Test Little (1988) proposed a multivariate test of Missing Completely at Random (MCAR) that tests for mean differences on every variable in the data set across subgroups that share the same missing data pattern by comparing the observed variable means for each pattern of missing data with the expected population means estimated using the expectation-maximization (EM) algorithm (i.e., EM maximum likelihood estimates). The test statistic is the sum of the squared standardized differences between the subsample means and the expected population means weighted by the estimated variance-covariance matrix and the number of observations within each subgroup (Enders, 2010). Under the null hypothesis that data are MCAR, the test statistic follows asymptotically a chi-square distribution with $\sum k_j - k$ degrees of freedom, where k_j is the number of complete variables for missing data pattern j , and k is the total number of variables. A statistically significant result provides evidence against MCAR.

Note that Little's MCAR test has a number of problems (see Enders, 2010).

- **First**, the test does not identify the specific variables that violates MCAR, i.e., the test does not identify potential correlates of missingness (i.e., auxiliary variables).
- **Second**, the test is based on multivariate normality, i.e., under departure from the normality assumption the test might be unreliable unless the sample size is large and is not suitable for categorical variables.
- **Third**, the test investigates mean differences assuming that the missing data pattern share a common covariance matrix, i.e., the test cannot detect covariance-based deviations from MCAR stemming from a Missing at Random (MAR) or Missing Not at Random (MNAR) mechanism because MAR and MNAR mechanisms can also produce missing data subgroups with equal means.
- **Fourth**, simulation studies suggest that Little's MCAR test suffers from low statistical power, particularly when the number of variables that violate MCAR is small, the relationship between the data and missingness is weak, or the data are MNAR (Thoemmes & Enders, 2007).
- **Fifth**, the test can only reject, but cannot prove the MCAR assumption, i.e., a statistically not significant result and failing to reject the null hypothesis of the MCAR test does not prove the null hypothesis that the data is MCAR.
- **Sixth**, under the null hypothesis the data are actually MCAR or MNAR, while a statistically significant result indicates that missing data are MAR or MNAR, i.e., MNAR cannot be ruled out regardless of the result of the test.

The function for performing Little's MCAR test is based on the `mlest` function from the `mvnmle` package which can handle up to 50 variables. Note that the `mcAR_test` function in

the **nanian** package is based on the `prelim.norm` function from the **norm** package. This function can handle about 30 variables, but with more than 30 variables specified in the argument `data`, the `prelim.norm` function might run into numerical problems leading to results that are not trustworthy (i.e., `p.value = 1`). In that case, the warning message `In norm::prelim.norm(data) : NAs introduced by coercion to integer range` is printed on the console.

Jamshidian and Jalal's Approach for Testing MCAR Jamshidian and Jalal (2010) proposed an approach for testing the Missing Completely at Random (MCAR) assumption based on two tests of multivariate normality and homogeneity of covariances among groups of cases with identical missing data patterns:

1. **In the first step**, missing data are multiply imputed ($m = 20$ times by default) using a non-parametric imputation method (`method = "npar"` by default) by Sirvastava and Dolatabadi (2009) or using a parametric imputation method assuming multivariate normality of data (`method = "normal"`) for each group of cases sharing a common missing data pattern.
2. **In the second step**, a modified Hawkins test for multivariate normality and homogeneity of covariances applicable to complete data consisting of groups with a small number of cases is performed. A statistically not significant result indicates no evidence against multivariate normality of data or homogeneity of covariances, while a statistically significant result provides evidence against multivariate normality of data or homogeneity of covariances (i.e., violation of the MCAR assumption). Note that the Hawkins test is a test of multivariate normality as well as homogeneity of covariance. Hence, a statistically significant test is ambiguous unless the researcher assumes multivariate normality of data.
3. **In the third step**, if the Hawkins test is statistically significant, the Anderson-Darling non-parametric test is performed. A statistically not significant result indicates evidence against multivariate normality of data but no evidence against homogeneity of covariances, while a statistically significant result provides evidence against homogeneity of covariances (i.e., violation of the MCAR assumption). However, no conclusions can be made about the multivariate normality of data when the Anderson-Darling non-parametric test is statistically significant.

In summary, a statistically significant result of both the Hawkins and the Anderson-Darling non-parametric test provides evidence against the MCAR assumption. The test statistic and the significance values of the Hawkins test and the Anderson-Darling non-parametric based on multiply imputed data sets are pooled by computing the median test statistic and significance value (`pool = "med"` by default) as suggested by Eekhout, Wiel, and Heymans (2017).

Note that out of the problems listed for the Little's MCAR test the first, second (i.e., approach is not suitable for categorical variables), fifth, and sixth problems also apply to the Jamshidian and Jalal's approach for testing the MCAR assumption.

In practice, rejecting or not rejecting the MCAR assumption may not be relevant as modern missing data handling methods like full information maximum likelihood (FIML) estimation, Bayesian estimation, or multiple imputation are asymptotically valid under the missing at random (MAR) assumption (Jamshidian & Yuan, 2014). It is more important to distinguish MAR from missing not at random (MNAR), but MAR and MNAR mechanisms cannot be distinguished without auxiliary information.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	matrix or data frame specified in data
args	specification of function arguments
result	list with result tables, i.e., little for the result table of the Little's MCAR test, jamjal for the list with results of the Jamshidian and Jalal's approach, hawkins for the result table of the Hawkins test, and anderson for the result table of the Anderson-Darling non-parametric test

Note

The code for Little's MCAR test is a modified copy of the LittleMCAR function in the **BaylorEdPsych** package by A. Alexander Beaujean. The code for Jamshidian and Jalal's approach is a modified copy of the TestMCARNormality function in the **MissMech** package by Mortaza Jamshidian, Siavash Jalal, Camden Jansen, and Mao Kobayashi (2024).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Beaujean, A. A. (2012). *BaylorEdPsych: R Package for Baylor University Educational Psychology Quantitative Courses*. R package version 0.5. <http://cran.nexr.com/web/packages/BaylorEdPsych/index.html>
- Eekhout, I., M. A. Wiel, & M. W. Heymans (2017). Methods for significance testing of categorical covariates in logistic regression models after multiple imputation: Power and applicability analysis. *BMC Medical Research Methodology*, 17:129. <https://doi.org/10.1186/s12874-017-0404-7>
- Enders, C. K. (2022). *Applied missing data analysis* (2nd ed.). The Guilford Press.
- Little, R. J. A. (1988). A test of Missing Completely at Random for multivariate data with missing values. *Journal of the American Statistical Association*, 83, 1198-1202. <https://doi.org/10.2307/2290157>
- Jamshidian, M., & Jalal, S. (2010). Tests of homoscedasticity, normality, and missing completely at random for incomplete multivariate data. *Psychometrika*, 75(4), 649-674. <https://doi.org/10.1007/s11336-010-9175-3>
- Jamshidian, M., & Yuan, K.H. (2014). Examining missing data mechanisms via homogeneity of parameters, homogeneity of distributions, and multivariate normality. *WIREs Computational Statistics*, 6(1), 56-73. <https://doi.org/10.1002/wics.1287>
- Mortaza, J., Siavash, J., Camden, J., & Kobayashi, M. (2024). *MissMech: Testing Homoscedasticity, Multivariate Normality, and Missing Completely at Random*. R package version 1.0.4. <https://doi.org/10.32614/CRAN.package>
- Srivastava, M.S., & Dolatabadi, M. (2009). Multiple imputation and other resampling scheme for imputing missing observations. *Journal of Multivariate Analysis*, 100, 1919-1937. <https://doi.org/10.1016/j.jmva.2009.06.001>
- Thoemmes, F., & Enders, C. K. (2007, April). *A structural equation model for testing whether data are missing completely at random*. Paper presented at the annual meeting of the American Educational Research Association, Chicago, IL.

See Also

[as.na](#), [na.as](#), [na.auxiliary](#), [na.coverage](#), [na.descript](#), [na.indicator](#), [na.pattern](#), [na.prop](#).

Examples

```
# Example 1: Perform Little's MCAR test and Jamshidian and Jalal's approach
na.test(airquality)

# Example 2: Perform Jamshidian and Jalal's approach
na.test(airquality, print = "jamjal")

## Not run:
# Example 3: Write results into a text file
na.test(airquality, write = "NA_Test.txt")

## End(Not run)
```

plot.misty.object	<i>Plots misty.object object</i>
-------------------	----------------------------------

Description

This function plots an `misty.object` object.

Usage

```
## S3 method for class 'misty.object'
plot(x, plot = x$args$plot, bar = x$args$bar,
     box = x$args$box, violin = x$args$violin, hist = x$args$hist,
     point = x$args$point, line = x$args$line, ci = x$args$ci,
     conf.level = x$args$conf.level, adjust = x$args$adjust,
     jitter = x$args$jitter, density = x$args$density,
     square = x$args$square, rotate = x$args$rotate,
     binwidth = x$args$binwidth, bins = x$args$bins,
     fill = x$args$fill, hist.alpha = x$args$hist.alpha,
     tile.alpha = x$args$tile.alpha, violin.alpha = x$args$violin.alpha,
     violin.trim = x$args$violin.trim, box.width = x$args$box.width,
     box.alpha = x$args$box.alpha, linetype = x$args$linetype,
     linewidth = x$args$linewidth, line.col = x$args$line.col,
     intercept = x$args$intercept, density.col = x$args$density.col,
     density.linewidth = x$args$density.linewidth,
     density.linetype = x$args$density.linetype,
     point.size = x$args$point.size, point.linewidth = x$args$point.linewidth,
     point.linetype = x$args$point.linetype,
     point.shape = x$args$point.shape, point.col = x$args$point.col,
     ci.col = x$args$ci.col, ci.linewidth = x$args$ci.linewidth,
     ci.linetype = x$args$ci.linetype, errorbar.width = x$args$errorbar.width,
     dodge.width = x$args$dodge.width, jitter.size = x$args$jitter.size,
```

```

jitter.width = x$args$jitter.width, jitter.height = x$args$jitter.height,
jitter.alpha = x$args$jitter.alpha, gray = x$args$gray,
start = x$args$start, end = x$args$end, color = x$args$color,
xlab = x$args$xlab, ylab = x$args$ylab,
xlim = x$args$xlim, ylim = x$args$ylim,
xbreaks = x$args$xbreaks, ybreaks = x$args$ybreaks,
axis.title.size = x$args$axis.title.sizes,
axis.text.size = x$args$axis.text.size,
strip.text.size = x$args$strip.text.size, title = x$args$title,
subtitle = x$args$subtitle, group.col = x$args$group.col,
plot.margin = x$args$plot.margin, legend.title = x$args$legend.title,
legend.position = x$args$legend.position,
legend.box.margin = x$args$legend.box.margin,
legend.key.size = x$args$legend.key.size,
legend.text.size = x$args$legend.text.size,
facet.ncol = x$args$facet.ncol, facet.nrow = x$args$facet.nrow,
facet.scales = x$args$facet.scales, filename = x$args$filename,
width = x$args$width, height = x$args$height, units = x$args$units,
dpi = x$args$dpi, check = TRUE, ...)

```

Arguments

<code>x</code>	misty.object object.
<code>plot</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>bar</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>box</code>	see 'Arguments' in the function <code>test.levene</code> .
<code>violin</code>	see 'Arguments' in the function <code>test.levene</code> .
<code>hist</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>point</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), or <code>test.welch</code> .
<code>line</code>	see 'Arguments' in the functions <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), or <code>test.t</code> .
<code>ci</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.t</code> , or <code>test.welch</code> .
<code>conf.level</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.t</code> , or <code>test.welch</code> .
<code>adjust</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>jitter</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>test.t</code> , <code>test.levene</code> , or <code>test.welch</code> .
<code>density</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>square</code>	see 'Arguments' in the function <code>na.pattern</code> .
<code>rotate</code>	see 'Arguments' in the function <code>na.pattern</code> .
<code>binwidth</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>bins</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>fill</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>hist.alpha</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).

tile.alpha	see 'Arguments' in the function na.pattern .
violin.alpha	see 'Arguments' in the function test.levene .
violin.trim	see 'Arguments' in the function test.levene .
box.width	see 'Arguments' in the function test.levene .
box.alpha	see 'Arguments' in the function test.levene .
linetype	see 'Arguments' in the functions ci.* (e.g., ci.cor) or test.t .
linewidth	see 'Arguments' in the function test.t .
line.col	see 'Arguments' in the functions ci.* (e.g., ci.cor).
intercept	see 'Arguments' in the functions ci.* (e.g., ci.cor).
density.col	see 'Arguments' in the functions ci.* (e.g., ci.cor).
density.linewidth	see 'Arguments' in the functions ci.* (e.g., ci.cor).
density.linetype	see 'Arguments' in the functions ci.* (e.g., ci.cor).
point.size	see 'Arguments' in the functions aov.b , aov.w , ci.* (e.g., ci.cor), test.t , or test.welch .
point.linewidth	see 'Arguments' in the functions ci.* (e.g., ci.cor).
point.linetype	see 'Arguments' in the functions ci.* (e.g., ci.cor).
point.shape	see 'Arguments' in the functions ci.* (e.g., ci.cor).
point.col	see 'Arguments' in the functions ci.* (e.g., ci.cor).
ci.col	see 'Arguments' in the functions ci.* (e.g., ci.cor).
ci.linewidth	see 'Arguments' in the functions ci.* (e.g., ci.cor).
ci.linetype	see 'Arguments' in the functions ci.* (e.g., ci.cor).
errorbar.width	see 'Arguments' in the functions aov.b , aov.w , ci.* (e.g., ci.cor), test.t , or test.welch .
dodge.width	see 'Arguments' in the functions ci.* (e.g., ci.cor).
jitter.size	see 'Arguments' in the functions aov.b , aov.w , test.levene , test.t , or test.welch .
jitter.width	see 'Arguments' in the functions aov.b , aov.w , test.levene , test.t , or test.welch .
jitter.height	see 'Arguments' in the functions aov.b , test.levene , test.t , or test.welch .
jitter.alpha	see 'Arguments' in the functions aov.b , aov.w , test.levene , test.t , or test.welch .
gray	see 'Arguments' in the functions multilevel.r2 or test.levene .
start	see 'Arguments' in the functions multilevel.r2 or test.levene .
end	see 'Arguments' in the functions multilevel.r2 or test.levene .
color	see 'Arguments' in the functions multilevel.r2 , na.pattern , or test.levene .
xlab	see 'Arguments' in the functions aov.b , aov.w , ci.* (e.g., ci.cor), test.levene , test.t , or test.welch .
ylab	see 'Arguments' in the functions aov.b , aov.w , ci.* (e.g., ci.cor), test.levene , test.t , or test.welch .

<code>xlim</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>ylim</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>xbreaks</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>ybreaks</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>axis.title.size</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>axis.text.size</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>strip.text.size</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>title</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>subtitle</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>group.col</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>plot.margin</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>legend.title</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>legend.position</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>legend.box.margin</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>) or <code>na.pattern</code> .
<code>legend.key.size</code>	see 'Arguments' in the function <code>na.pattern</code> .
<code>legend.text.size</code>	see 'Arguments' in the functions <code>na.pattern</code> .
<code>facet.ncol</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>facet.nrow</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>facet.scales</code>	see 'Arguments' in the functions <code>ci.*</code> (e.g., <code>ci.cor</code>).
<code>filename</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>width</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>height</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>units</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>dpi</code>	see 'Arguments' in the functions <code>aov.b</code> , <code>aov.w</code> , <code>ci.*</code> (e.g., <code>ci.cor</code>), <code>test.levene</code> , <code>test.t</code> , or <code>test.welch</code> .
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>...</code>	further arguments passed to or from other methods.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

print.misty.object *Print misty.object object*

Description

This function prints an `misty.object` object.

Usage

```
## S3 method for class 'misty.object'
print(x,
      print = x$args$print, tri = x$args$tri, freq = x$args$freq,
      hypo = x$args$hypo, descript = x$args$descript, epsilon = x$args$epsilon,
      effsize = x$args$effsize, posthoc = x$args$posthoc, split = x$args$split,
      table = x$args$table, digits = x$args$digits, p.digits = x$args$p.digits,
      icc.digits = x$args$icc.digits, r.digits = x$args$r.digits,
      ess.digits = x$args$ess.digits, mcse.digits = x$args$mcse.digits,
      sort.var = x$args$sort.var, order = x$args$order, check = TRUE, ...)
```

Arguments

<code>x</code>	<code>misty.object</code> object.
<code>print</code>	a character string or character vector indicating which results to to be printed on the console.
<code>tri</code>	a character string or character vector indicating which triangular of the matrix to show on the console, i.e., both for upper and lower triangular, lower for the lower triangular, and upper for the upper triangular.
<code>freq</code>	logical: if TRUE, absolute frequencies will be included in the cross tabulation (<code>crosstab()</code> function).
<code>hypo</code>	logical: if TRUE, null and alternative hypothesis are shown on the console (test.t , test.welch , test.z function).
<code>descript</code>	logical: if TRUE, descriptive statistics are shown on the console (test.t , test.welch , test.z function).
<code>epsilon</code>	logical: if TRUE, box indices of sphericity (epsilon) are shown on the console (aov.w).
<code>effsize</code>	logical: if TRUE, effect size measure(s) is shown on the console (test.t , test.welch , test.z function). test.z function).
<code>posthoc</code>	logical: if TRUE, post hoc test for multiple comparison is shown on the console (test.welch).
<code>split</code>	logical: if TRUE, output table is split by variables when specifying more than one variable in <code>x</code> (freq).

table	logical: if TRUE, a frequency table with number of observed values ("nOb"), percent of observed values ("pOb"), number of missing values ("nNA"), and percent of missing values ("pNA") is printed for each variable on the console (na.descript() function).
digits	an integer value indicating the number of decimal places to be used for displaying results.
p.digits	an integer indicating the number of decimal places to be used for displaying <i>p</i> -values.
icc.digits	an integer indicating the number of decimal places to be used for displaying intraclass correlation coefficients (multilevel.descript() and multilevel.icc() function).
r.digits	an integer value indicating the number of decimal places to be used for displaying R-hat values.
ess.digits	an integer value indicating the number of decimal places to be used for displaying effective sample sizes.
mcse.digits	an integer value indicating the number of decimal places to be used for displaying monte carlo standard errors.
sort.var	logical: if TRUE, output is sorted by variables.
order	logical: if TRUE, variables are ordered from left to right in increasing order of missing values (na.descript() function).
check	logical: if TRUE, argument specification is checked.
...	further arguments passed to or from other methods.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

read.data

Read Data File in Table Format, SPSS, Excel, or Stata DTA File

Description

This function reads a (1) data file in CSV (.csv), DAT (.dat), or TXT (.txt) format using the fread function from the **data.table** package, (2) SPSS file (.sav) using the read.sav function, (3) Excel file (.xlsx) using the read.xlsx function, or a (4) Stata DTA file (.dta) using the read.dta function in the **misty** package.

Usage

```
read.data(file, sheet = NULL, header = TRUE, select = NULL, drop = NULL,
          sep = "auto", dec = "auto", use.value.labels = FALSE,
          use.missings = TRUE, na.strings = c("NA", ""),
          stringsAsFactors = FALSE, formats = FALSE, label = FALSE,
          labels = FALSE, missing = FALSE, widths = FALSE, as.data.frame = TRUE,
          encoding = c("unknown", "UTF-8", "Latin-1"), check = TRUE)
```

Arguments

file	a character string indicating the name of the data file with the file extension .csv, .dat, .txt, .sav, .xlsx, or .dta. Note that the function will select an appropriate read-function depending on the file extension.
sheet	a character string indicating the name of a Excel sheet or a numeric value indicating the position of the Excel sheet to read. By default the first sheet will be read when reading an Excel file (.xlsx).
header	logical: if TRUE (default), the first row is used as column names when reading an Excel file (.xlsx), if FALSE default names are used. A character vector giving a name for each column can also be used.
select	a character vector of column names or numeric vector to keep, drop the rest. See the help page of the fread function in the data.table package.
drop	a character vector of column names or numeric vector to drop, keep the rest.
sep	a character string indicating the separator between columns for the fread function when reading data in CSV (.csv), DAT (.dat), or TXT (.txt) format.
dec	a character string indicating the decimal separator for the fread function when reading data in CSV (.csv), DAT (.dat), or TXT (.txt) format.
use.value.labels	logical: if TRUE, variables with value labels are converted into factors.
use.missings	logical: if TRUE (default), user-defined missing values are converted into NAs.
na.strings	a character vector of strings which are to be interpreted as NA values.
stringsAsFactors	logical: if TRUE, character vectors are converted to factors.
formats	logical: if TRUE, variable formats are shown in an attribute for all variables.
label	logical: if TRUE, variable labels are shown in an attribute for all variables.
labels	logical: if TRUE, value labels are shown in an attribute for all variables.
missing	logical: if TRUE, value labels for user-defined missings are shown in an attribute for all variables.
widths	logical: if TRUE, widths are shown in an attribute for all variables.
as.data.frame	logical: if TRUE (default), function returns a regular data frame; if FALSE function returns a tibble or data.table.
encoding	a character string indicating the encoding, i.e., "unknown", "UTF-8", or "Latin-1" (default).
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a data frame, tibble, or data table.

Author(s)

Takuya Yanagida

References

Barrett, T., Dowle, M., Srinivasan, A., Gorecki, J., Chirico, M., Hocking, T., & Schwendinger, B. (2024). data.table: Extension of 'data.frame'. R package version 1.16.0. <https://CRAN.R-project.org/package=data.table>

Wickham H, Miller E, Smith D (2023). haven: Import and Export 'SPSS', 'Stata' and 'SAS' Files. R package version 2.5.3. <https://CRAN.R-project.org/package=haven>

See Also

[write.data](#), [read.sav](#), [write.sav](#), [write.xlsx](#), [read.dta](#), [write.dta](#), [read.mplus](#), [write.mplus](#)

Examples

```
## Not run:

# Example 1: Read CSV data file
dat <- read.data("CSV_Data.csv")

# Example 2: Read DAT data file
dat <- read.data("DAT_Data.dat")

# Example 3: Read TXT data file
dat <- read.data("TXT_Data.txt")

# Example 4: Read SPSS data file
dat <- read.data("SPSS_Data.sav")

# Example 5: Read Excel data file
dat <- read.data("Excel_Data.xlsx")

# Example 6: Read Stata data file
dat <- read.data("Stata_Data.dta")

## End(Not run)
```

read.dta	<i>Read Stata DTA File</i>
----------	----------------------------

Description

This function calls the read_dta function in the **haven** package by Hadley Wickham, Evan Miller and Danny Smith (2023) to read a Stata DTA file.

Usage

```
read.dta(file, use.value.labels = FALSE, formats = FALSE, label = FALSE, labels = FALSE,
         missing = FALSE, widths = FALSE, as.data.frame = TRUE, check = TRUE)
```

Arguments

file	a character string indicating the name of the Stata data file with or without file extension '.dta', e.g., "Stata_Data.dta" or "Stata_Data".
use.value.labels	logical: if TRUE, variables with value labels are converted into factors.
formats	logical: if TRUE (default), variable formats are shown in an attribute for all variables.
label	logical: if TRUE, variable labels are shown in an attribute for all variables.
labels	logical: if TRUE, value labels are shown in an attribute for all variables.
missing	logical: if TRUE, convert tagged missing values to regular R NA.
widths	logical: if TRUE, widths are shown in an attribute for all variables.
as.data.frame	logical: if TRUE (default), function returns a regular data frame; if FALSE function returns a tibble.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a data frame or tibble.

Note

This function is a modified copy of the `read_dta()` function in the **haven** package by Hadley Wickham, Evan Miller and Danny Smith (2023).

Author(s)

Hadley Wickham and Evan Miller

References

Wickham H, Miller E, Smith D (2023). *haven: Import and Export 'SPSS', 'Stata' and 'SAS' Files*. R package version 2.5.3. <https://CRAN.R-project.org/package=haven>

See Also

[read.data](#), [write.data](#), [read.sav](#), [write.sav](#), [write.xlsx](#), [write.dta](#), [read.mplus](#), [write.mplus](#)

Examples

```
## Not run:

read.dta("Stata_Data.dta")
read.dta("Stata_Data")

# Example 2: Read Stata data, convert variables with value labels into factors
read.dta("Stata_Data.dta", use.value.labels = TRUE)

# Example 3: Read Stata data as tibble
```

```
read.dta("Stata_Data.dta", as.data.frame = FALSE)

## End(Not run)
```

read.mplus

Read Mplus Data File and Variable Names

Description

This function reads a Mplus data file and/or Mplus input/output file to return a data frame with variable names extracted from the Mplus input/output file. Note that by default -99 in the Mplus data file is replaced with to NA.

Usage

```
read.mplus(file, sep = "", input = NULL, na = -99, print = FALSE, return.var = FALSE,
           encoding = "UTF-8-BOM", check = TRUE)
```

Arguments

file	a character string indicating the name of the Mplus data file with or without the file extension .dat, e.g., "Mplus_Data.dat" or "Mplus_Data". Note that it is not necessary to specify this argument when return.var = TRUE.
sep	a character string indicating the field separator (i.e., delimiter) used in the data file specified in file. By default, the separator is 'white space', i.e., one or more spaces, tabs, newlines or carriage returns.
input	a character string indicating the Mplus input (.inp) or output file (.out) in which the variable names are specified in the VARIABLE: section. Note that if input = NULL, this function is equivalent to read.table(file).
na	a numeric vector indicating values to replace with NA. By default, -99 is replaced with NA. If -99 is not a missing value change the argument to NULL.
print	logical: if TRUE, variable names are printed on the console.
return.var	logical: if TRUE, the function returns the variable names extracted from the Mplus input or output file only.
encoding	character string declaring the encoding used on file so the character data can be re-encoded. See the 'Encoding' section of the help page for the file function, the 'R Data Import/Export Manual' and 'Note'.
check	logical: if TRUE (default), argument specification is checked.

Value

A data frame containing a representation of the data in the file.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Muthen, L. K., & Muthen, B. O. (1998-2017). *Mplus User's Guide* (8th ed.). Muthen & Muthen.

See Also

[read.data](#), [write.data](#), [read.sav](#), [write.sav](#), [write.xlsx](#), [read.dta](#), [write.dta](#), [write.mplus](#)

Examples

```
## Not run:

# Example 1: Read Mplus data file and variable names extracted from the Mplus input file
dat <- read.mplus("Mplus_Data.dat", input = "Mplus_Input.inp")

# Example 2: Read Mplus data file and variable names extracted from the Mplus input file,
# print variable names on the console
dat <- read.mplus("Mplus_Data.dat", input = "Mplus_Input.inp", print = TRUE)

# Example 3: Read variable names extracted from the Mplus input file
varnames <- read.mplus(input = "Mplus_Input.inp", return.var = TRUE)

## End(Not run)
```

read.sav

Read SPSS File

Description

This function calls the `read_spss` function in the **haven** package by Hadley Wickham, Evan Miller and Danny Smith (2023) to read an SPSS file.

Usage

```
read.sav(file, use.value.labels = FALSE, use.missings = TRUE, formats = FALSE,
         label = FALSE, labels = FALSE, missing = FALSE, widths = FALSE,
         as.data.frame = TRUE, check = TRUE)
```

Arguments

<code>file</code>	a character string indicating the name of the SPSS data file with or without file extension '.sav', e.g., "SPSS_Data.sav" or "SPSS_Data".
<code>use.value.labels</code>	logical: if TRUE, variables with value labels are converted into factors.
<code>use.missings</code>	logical: if TRUE (default), user-defined missing values are converted into NAs.
<code>formats</code>	logical: if TRUE, variable formats are shown in an attribute for all variables.
<code>label</code>	logical: if TRUE, variable labels are shown in an attribute for all variables.
<code>labels</code>	logical: if TRUE, value labels are shown in an attribute for all variables.

missing	logical: if TRUE, value labels for user-defined missings are shown in an attribute for all variables.
widths	logical: if TRUE, widths are shown in an attribute for all variables.
as.data.frame	logical: if TRUE (default), function returns a regular data frame; if FALSE function returns a tibble.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a data frame or tibble.

Author(s)

Hadley Wickham, Evan Miller and Danny Smith

References

Wickham H, Miller E, & Smith D (2023). *haven: Import and Export 'SPSS', 'Stata' and 'SAS' Files*. R package version 2.5.3. <https://CRAN.R-project.org/package=haven>

See Also

[read.data](#), [write.data](#), [write.sav](#), [write.xlsx](#), [read.dta](#), [write.dta](#), [read.mplus](#), [write.mplus](#)

Examples

```
## Not run:

# Example 1: Read SPSS data file
read.sav("SPSS_Data.sav")
read.sav("SPSS_Data")

# Example 2: Read SPSS data file, convert variables with value labels into factors
read.sav("SPSS_Data.sav", use.value.labels = TRUE)

# Example 3: Read SPSS data file, user-defined missing values are not converted into NAs
read.sav("SPSS_Data.sav", use.missing = FALSE)

# Example 4: Read SPSS data file as tibble
read.sav("SPSS_Data.sav", as.data.frame = FALSE)

## End(Not run)
```

read.xlsx

*Read Excel File***Description**

This function calls the `read_xlsx()` function in the **readxl** package by Hadley Wickham and Jennifer Bryan (2019) to read an Excel file (.xlsx).

Usage

```
read.xlsx(file, sheet = NULL, header = TRUE, range = NULL,
  coltypes = c("skip", "guess", "logical", "numeric", "date", "text", "list"),
  na = "", trim = TRUE, skip = 0, nmax = Inf, guessmax = min(1000, nmax),
  progress = readxl::readxl_progress(), name.repair = "unique",
  as.data.frame = TRUE, check = TRUE)
```

Arguments

file	a character string indicating the name of the Excel data file with or without file extension '.xlsx', e.g., "My_Excel_Data.xlsx" or "My_Excel_Data".
sheet	a character string indicating the name of a sheet or a numeric value indicating the position of the sheet to read. By default the first sheet will be read.
header	logical: if TRUE (default), the first row is used as column names, if FALSE default names are used. A character vector giving a name for each column can also be used. If coltypes as a vector is provided, colnames can have one entry per column, i.e. have the same length as coltypes, or one entry per unskipped column.
range	a character string indicating the cell range to read from, e.g. typical Excel ranges like "B3:D87", possibly including the sheet name like "Data!B2:G14". Interpreted strictly, even if the range forces the inclusion of leading or trailing empty rows or columns. Takes precedence over skip, nmax and sheet.
coltypes	a character vector containing one entry per column from these options "skip", "guess", "logical", "numeric", "date", "text" or "list". If exactly one coltype is specified, it will be recycled. By default (i.e., coltypes = NULL) coltypes will be guessed. The content of a cell in a skipped column is never read and that column will not appear in the data frame output. A list cell loads a column as a list of length 1 vectors, which are typed using the type guessing logic from coltypes = NULL, but on a cell-by-cell basis.
na	a character vector indicating strings to interpret as missing values. By default, blank cells will be treated as missing data.
trim	logical: if TRUE (default), leading and trailing whitespace will be trimmed.
skip	a numeric value indicating the minimum number of rows to skip before reading anything, be it column names or data. Leading empty rows are automatically skipped, so this is a lower bound. Ignored if the argument range is specified.

nmax	a numeric value indicating the maximum number of data rows to read. Trailing empty rows are automatically skipped, so this is an upper bound on the number of rows in the returned data frame. Ignored if the argument range is specified.
guessmax	a numeric value indicating the maximum number of data rows to use for guessing column types.
progress	display a progress spinner? By default, the spinner appears only in an interactive session, outside the context of knitting a document, and when the call is likely to run for several seconds or more.
name.repair	a character string indicating the handling of column names. By default, the function ensures column names are not empty and are unique.
as.data.frame	logical: if TRUE (default), function returns a regular data frame; if FALSE function returns a tibble.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a data frame or tibble.

Author(s)

Hadley Wickham and Jennifer Bryan

References

Wickham H, Miller E, Smith D (2023). *readxl: Read Excel Files*. R package version 1.4.3. <https://CRAN.R-project.org/package=readxl>

See Also

[read.data](#), [write.data](#), [read.sav](#), [write.sav](#), [write.xlsx](#), [read.dta](#), [write.dta](#), [read.mplus](#), [write.mplus](#)

Examples

```
## Not run:

# Example 1: Read Excel file (.xlsx)
read.xlsx("data.xlsx")

# Example 1: Read Excel file (.xlsx), use default names as column names
read.xlsx("data.xlsx", header = FALSE)

# Example 2: Read Excel file (.xlsx), interpret -99 as missing values
read.xlsx("data.xlsx", na = "-99")

# Example 3: Read Excel file (.xlsx), use x1, x2, and x3 as column names
read.xlsx("data.xlsx", header = c("x1", "x2", "x3"))

# Example 4: Read Excel file (.xlsx), read cells A1:B5
```

```

read.xlsx("data.xlsx", range = "A1:B5")

# Example 5: Read Excel file (.xlsx), skip 2 rows before reading data
read.xlsx("data.xlsx", skip = 2)

# Example 5: Read Excel file (.xlsx), return a tibble
read.xlsx("data.xlsx", as.data.frame = FALSE)

## End(Not run)

```

rec

*Recode Variable***Description**

This function recodes numeric vectors, character vectors, or factors according to recode specifications.

Usage

```

rec(data, ..., spec, as.factor = FALSE, levels = NULL, append = TRUE,
     name = ".e", as.na = NULL, table = FALSE, check = TRUE)

```

Arguments

<code>data</code>	a numeric vector, character vector, factor, or data frame.
<code>...</code>	an expression indicating the variable names in data, e.g., <code>rec(dat, x1, x2, x3, spec = "1 = 0")</code>). Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>spec</code>	a character string of recode specifications (see 'Details').
<code>as.factor</code>	logical: if TRUE, character vector will be coerced to a factor.
<code>levels</code>	a character vector for specifying the levels in the returned factor.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>append</code>	logical: if TRUE (default), centered variable(s) are appended to the data frame specified in the argument data.
<code>name</code>	a character string or character vector indicating the names of the recoded variables. By default, variables are named with the ending <code>".r"</code> Resulting in e.g. <code>"x1.r"</code> and <code>"x2.r"</code> . Variable names can also be specified using a character vector matching the number of variables specified in data (e.g., <code>name = c("recode.x1", "recode.x2")</code>).
<code>table</code>	logical: if TRUE, a cross table variable x recoded variable is printed on the console if only one variable is specified in data.
<code>check</code>	logical: if TRUE (default), argument specification is checked.

Details

Recode specifications appear in a character string, separated by semicolons (see the examples below), of the form `input = output`. If an input value satisfies more than one specification, then the first (from left to right) applies. If no specification is satisfied, then the input value is carried over to the result. NA is allowed in input and output. Several recode specifications are supported:

Single Value For example, `spec = "0 = NA"`.

Vector of Values For example, `spec = "c(7, 8, 9) = 'high'"`.

Range of Values For example, `spec = "7:9 = 'C'"`. The special values `lo` (lowest value) and `hi` (highest value) may appear in a range. For example, `spec = "lo:10 = 1"`. Note that `:` is not the R sequence operator. In addition you may not use `:` with the collect operator, e.g., `spec = "c(1, 3, 5:7)"` will cause an error.

else For example, `spec = "0 = 1; else = NA"`. Everything that does not fit a previous specification. Note that `else` matches all otherwise unspecified values on input, including NA.

Value

Returns a numeric vector or data frame with the same length or same number of rows as data containing the recoded coded variable(s).

Note

This function was adapted from the `recode()` function in the **car** package by John Fox and Sanford Weisberg (2019).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Fox, J., & Weisberg S. (2019). *An R Companion to Applied Regression* (3rd ed.). Thousand Oaks CA: Sage. URL: <https://socialsciences.mcmaster.ca/jfox/Books/Companion/>

See Also

[coding](#), [item.reverse](#)

Examples

```
#-----
# Example 1: Numeric vector
x.num <- c(1, 2, 4, 5, 6, 8, 12, 15, 19, 20)

# Example 1a: Recode 5 = 50 and 19 = 190
rec(x.num, spec = "5 = 50; 19 = 190")

# Example 1b: Recode 1, 2, and 5 = 100 and 4, 6, and 7 = 200 and else = 300
rec(x.num, spec = "c(1, 2, 5) = 100; c(4, 6, 7) = 200; else = 300")
```

```

# Example 1c: Recode lowest value to 10 = 100 and 11 to highest value = 200
rec(x.num, spec = "lo:10 = 100; 11:hi = 200")

# Example 1d: Recode 5 = 50 and 19 = 190 and check recoding
rec(x.num, spec = "5 = 50; 19 = 190", table = TRUE)

#-----
# Example 2: Character vector
x.chr <- c("a", "c", "f", "j", "k")

# Example 2a: Recode a to x
rec(x.chr, spec = "'a' = 'X'")

# Example 2b: Recode a and f to x, c and j to y, and else to z
rec(x.chr, spec = "c('a', 'f') = 'x'; c('c', 'j') = 'y'; else = 'z'")

# Example 2c: Recode a to x and coerce to a factor
rec(x.chr, spec = "'a' = 'X'", as.factor = TRUE)

#-----
# Example 3: Factor
x.fac <- factor(c("a", "b", "a", "c", "d", "d", "b", "b", "a"))

# Example 3a: Recode a to x, factor levels ordered alphabetically
rec(x.fac, spec = "'a' = 'x'")

# Example 3b: Recode a to x, user-defined factor levels
rec(x.fac, spec = "'a' = 'x'", levels = c("x", "b", "c", "d"))

#-----
# Example 4: Multiple variables
dat <- data.frame(x1.num = c(1, 2, 4, 5, 6),
                  x2.num = c(5, 19, 2, 6, 3),
                  x1.chr = c("a", "c", "f", "j", "k"),
                  x2.chr = c("b", "c", "a", "d", "k"),
                  x1.fac = factor(c("a", "b", "a", "c", "d")),
                  x2.fac = factor(c("b", "a", "d", "c", "e")))

# Example 4a: Recode numeric vector and attach to 'dat'
cbind(dat, rec(dat[, c("x1.num", "x2.num")], spec = "5 = 50; 19 = 190"))

# Alternative specification using the 'data' argument,
rec(dat, x1.num, x2.num, spec = "5 = 50; 19 = 190")

# Example 4b: Recode character vector and attach to 'dat'
cbind(dat, rec(dat[, c("x1.chr", "x2.chr")], spec = "'a' = 'X'"))

# Example 4c: Recode factor vector and attach to 'dat'
cbind(dat, rec(dat[, c("x1.fac", "x2.fac")], spec = "'a' = 'X'"))

```

restart	<i>Restart R Session</i>
---------	--------------------------

Description

This function restarts the RStudio session and is equivalent to using the menu item Session - Restart R.

Usage

restart()

Details

The function call `executeCommand("restartR")` in the package **rstudioapi** is used to restart the R session. Note that the function `restartSession()` in the package **rstudioapi** is not equivalent to the menu item Session - Restart R since it does not unload packages loaded during an R session.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Ushey, K., Allaire, J., Wickham, H., & Ritchie, G. (2022). rstudioapi: Safely access the RStudio API. R package version 0.14. <https://CRAN.R-project.org/package=rstudioapi>

Examples

```
## Not run:

# Example 1: Restart the R Session
restart()

## End(Not run)
```

result.lca	<i>Summary Result Table and Grouped Bar Charts for Latent Class Analysis Estimated in Mplus</i>
------------	---

Description

This function reads all Mplus output files from latent class analysis in subfolders to create a summary result table and bar charts for each latent class solution separately. By default, the function reads output files in all subfolders of the current working directory. Optionally, bar charts for each latent class solution can be requested by setting the argument `plot` to TRUE. Note that subfolders with only one Mplus output file are excluded.

Usage

```
result.lca(folder = getwd(), exclude = NULL, sort.n = TRUE, sort.p = TRUE,
  plot = FALSE, group.ind = TRUE, ci = TRUE, conf.level = 0.95, adjust = TRUE,
  axis.title = 7, axis.text = 7, levels = NULL, labels = NULL,
  ylim = NULL, ylab = "Mean Value", breaks = ggplot2::waiver(),
  errorbar.width = 0.1, legend.title = 7, legend.text = 7, legend.key.size = 0.4,
  gray = FALSE, start = 0.15, end = 0.85, dpi = 600,
  width = "n.ind", height = 4, digits = 1, p.digits = 3,
  write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

folder	a character vector indicating the name of the subfolders to be excluded from the summary result table.
exclude	a character vector indicating the name of the subfolders excluded from the result tables.
sort.n	logical: if TRUE (default), result table is sorted according to the number of classes within each folder.
sort.p	logical: if TRUE (default), class proportions are sorted decreasing.
plot	logical: if TRUE, bar charts with error bars for confidence intervals are saved in the folder <code>_Plots</code> within subfolders. Note that plots are only available for LCA with continuous or count indicator variables.
group.ind	logical: if TRUE (default), latent class indicators are represented by separate bars, if FALSE latent classes are represented by separate bars.
ci	logical: if TRUE (default), confidence intervals are added to the bar charts.
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
adjust	logical: if TRUE (default), difference-adjustment for the confidence intervals is applied.
axis.title	a numeric value specifying the size of the axis title.
axis.text	a numeric value specifying the size of the axis text
levels	a character string specifying the order of the indicator variables shown on the x-axis.
labels	a character string specifying the labels of the indicator variables shown on the x-axis.
ylim	a numeric vector of length two specifying limits of the y-axis.
ylab	a character string specifying the label of the y-axis.
breaks	a numeric vector specifying the points at which tick-marks are drawn at the y-axis.
errorbar.width	a numeric vector specifying the width of the error bars. By default, the width of the error bars is 0.1 plus number of classes divided by 30.
legend.title	a numeric value specifying the size of the legend title.
legend.text	a numeric value specifying the size of the legend text.

legend.key.size	a numeric value specifying the size of the legend keys.
gray	logical: if TRUE, bar charts are drawn in gray scale.
start	a numeric value between 0 and 1 specifying the gray value at the low end of the palette.
end	a numeric value between 0 and 1 specifying the gray value at the high end of the palette.
dpi	a numeric value specifying the plot resolution when saving the bar chart.
width	a numeric value specifying the width of the plot when saving the bar chart. By default, the width is number of indicators plus number of classes divided by 2.
height	a numeric value specifying the height of the plot when saving the bar chart.
digits	an integer value indicating the number of decimal places to be used for displaying results. Note that the scaling correction factor is displayed with digits plus 1 decimal places.
p.digits	an integer value indicating the number of decimal places to be used for displaying <i>p</i> -values, entropy value, and class proportions.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown.

Details

The result summary table comprises following entries:

- "Folder": Subfolder from which the group of Mplus outputs files were summarized.
- "#Class": Number of classes (i.e., CLASSES ARE c(#Class)).
- "Conv": Model converged, TRUE or FALSE (i.e., THE MODEL ESTIMATION TERMINATED NORMALLY).
- "#Param": Number of estimated parameters (i.e., Number of Free Parameters).
- "logLik": Log-likelihood of the estimated model (i.e., H0 Value).
- "Scale": Scaling correction factor (i.e., H0 Scaling Correction Factor for). Provided only when ESTIMATOR IS MLR.
- "LL Rep": Best log-likelihood replicated, TRUE or FALSE (i.e., THE BEST LOGLIKELIHOOD VALUE HAS BEEN REPLICATED).
- "AIC": Akaike information criterion (i.e., Akaike (AIC)).
- "CAIC": Consistent AIC, not reported in the Mplus output, but simply BIC + #Param.
- "BIC": Bayesian information criterion (i.e., Bayesian (BIC)).
- "Chi-Pear": Pearson chi-square test of model fit (i.e., Pearson Chi-Square), only available when indicators are count or ordered categorical.

- "Chi-LRT": Likelihood ratio chi-square test of model fit (i.e., Likelihood Ratio Chi-Square), only available when indicators are count or ordered categorical.
- "SABIC": Sample-size adjusted BIC (i.e., Sample-Size Adjusted BIC).
- "LMR-LRT": Significance value (p -value) of the Vuong-Lo-Mendell-Rubin test (i.e., VUONG-LO-MENDELL-RUBIN LIKELIHOOD RATIO TEST). Provided only when OUTPUT: TECH11.
- "A-LRT": Significance value (p -value) of the Adjusted Lo-Mendell-Rubin Test (i.e., LO-MENDELL-RUBIN ADJUSTED LRT TEST). Provided only when OUTPUT: TECH11.
- "BLRT": Significance value (p -value) of the bootstrapped likelihood ratio test. Provided only when OUTPUT: TECH14.
- "Entropy": Sample-size adjusted BIC (i.e., Entropy).
- "p1": Class proportion of the first class based on the estimated posterior probabilities (i.e., FINAL CLASS COUNTS AND PROPORTIONS).
- "p2": Class proportion of the second class based on the estimated posterior probabilities (i.e., FINAL CLASS COUNTS AND PROPORTIONS).

Value

Returns an object, which is a list with following entries:

call	function call
type	type of analysis
output	list with all Mplus outputs
args	specification of function arguments
result	list with result tables, i.e., summary for the summary result table, mean_var for the result table with means and variances for each latent class separately, mean for the result table with means for each latent class separately, and var for the result table with variances for each latent class separately

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Masyn, K. E. (2013). Latent class analysis and finite mixture modeling. In T. D. Little (Ed.), *The Oxford handbook of quantitative methods: Statistical analysis* (pp. 551–611). Oxford University Press.
- Muthen, L. K., & Muthen, B. O. (1998-2017). *Mplus User's Guide* (8th ed.). Muthen & Muthen.

See Also

[mplus.lca](#), [mplus.run](#), [read.mplus](#), [write.mplus](#)

Examples

```
## Not run:

# Load data set "HolzingerSwineford1939" in the lavaan package
data("HolzingerSwineford1939", package = "lavaan")

# Run LCA with k = 1 to k = 6 classes
mplus.lca(HolzingerSwineford1939, ind = c("x1", "x2", "x3", "x4"),
          mplus.run = TRUE)

# Example 1a: Read Mplus output files, create result table, write table, and save plots
result.lca(write = "LCA.xlsx", plot = TRUE)

# Example 1b: Write results into a text file
result.lca(write = "LCA.txt")

#-----
# Example 2: Draw bar chart manually

library(ggplot2)

# Collect LCA results
lca.result <- result.lca()

# Result table with means
means <- lca.result$result$mean

# Extract results from variance-covariance structure A with 4 latent classes
plotdat <- means[means$folder == "A_Invariant-Theta_Diagonal-Sigma" & means$nclass == 4, ]

# Draw bar chart
ggplot(plotdat, aes(ind, est, group = class, fill = class)) +
  geom_bar(stat = "identity", position = "dodge", color = "black",
          linewidth = 0.1) +
  geom_errorbar(aes(ymin = low, ymax = upp), width = 0.23,
               linewidth = 0.2, position = position_dodge(0.9)) +
  scale_x_discrete("") +
  scale_y_continuous("Mean Value", limits = c(0, 9),
                     breaks = seq(0, 9, by = 1)) +
  labs(fill = "Latent Class") +
  guides(fill = guide_legend(nrow = 1L)) +
  theme(axis.title = element_text(size = 11),
        axis.text = element_text(size = 11),
        legend.position = "bottom",
        legend.key.size = unit(0.5, 'cm'),
        legend.title = element_text(size = 11),
        legend.text = element_text(size = 11),
        legend.box.spacing = unit(-9L, "pt"))

# Save bar chart
ggsave("LCA_4-Class.png", dpi = 600, width = 6, height = 4)
```

```
## End(Not run)
```

robust.lmer

Robust Estimation of Multilevel and Linear Mixed-Effects Models

Description

This function estimates a multilevel and linear mixed-effects model based on a robust estimation method using the `rlmer()` function from the **robustlmm** package that down-weights observations depending on robustness weights computed by robustification of the scoring equations and an application of the Design Adaptive Scale (DAS) approach.

Usage

```
robust.lmer(model, method = c("DAStau", "DASvar"), setting = c("RSEn", "RSEa"),
  digits = 2, p.digits = 3, write = NULL, append = TRUE, check = TRUE,
  output = TRUE)
```

Arguments

<code>model</code>	a fitted model of class <code>"lmerMod"</code> or <code>"lmerModLmerTest"</code> .
<code>method</code>	a character string indicating the method used for estimating theta and sigma, i.e., <code>"DAStau"</code> (default) for using numerical quadrature for computing the consistency factors and <code>"DASvar"</code> for computing the consistency factors using a direct approximation. Note that <code>"DAStau"</code> is slower than <code>"DASvar"</code> but yields more accurate results. However, complex models with correlated random effects with more than one correlation term can only be estimated using <code>"DASvar"</code> . See help page of the <code>rlmer()</code> function in the R package <code>robustlmm</code> for more details.
<code>setting</code>	a character string indicating the setting for the parameters used for computing the robustness weights, i.e., <code>"RSEn"</code> (default) and <code>"RSEa"</code> for higher asymptotic efficiency which results in lower robustness. See help page of the <code>rlmer()</code> function in the R package <code>robustlmm</code> for more details.
<code>digits</code>	an integer value indicating the number of decimal places to be used.
<code>p.digits</code>	an integer value indicating the number of decimal places to be used for displaying <i>p</i> -value.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension <code>".txt"</code> (e.g., <code>"Output.txt"</code>) or Excel file with file extension <code>".xlsx"</code> (e.g., <code>"Output.xlsx"</code>). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if <code>TRUE</code> (default), output will be appended to an existing text file with extension <code>".txt"</code> specified in <code>write</code> , if <code>FALSE</code> existing text file will be overwritten.
<code>check</code>	logical: if <code>TRUE</code> (default), argument specification is checked.
<code>output</code>	logical: if <code>TRUE</code> (default), output is shown on the console.

Details

Function specification and Function Arguments The function `rlmer` from the **robustlmm** package is specified much like the function `lmer` from the `lme4` package, i.e., a formula object and data frame is specified as the first and second argument. However, the `robust.lmer` function requires a fitted `"lmerMod"` or `"lmerModLmerTest"` that is used to re-estimate the model using the robust estimation method. Note that the function `rlmer` provides the additional arguments `rho.e`, `rho.b`, `rho.sigma.e`, `rho.sigma.b`, `rel.tol`, `max.iter`, `verbose`, `doFit`, `init`, and `initTheta` that are not supported by the `robust.lmer` function. See help page of the `rlmer()` function in the R package `robustlmm` for more details.

Statistical Significance Testing The function `rlmer` from the **robustlmm** package does not provide any degrees of freedom or significance values. When specifying a `"lmerModLmerTest"` object for the argument `model`, the `robust.lmer` function uses the Satterthwaite or Kenward-Roger degrees of freedom from the `"lmerModLmerTest"` object to compute significance values for the regression coefficients based on parameter estimates and standard error of the robust multilevel mixed-effects (see Sleegers et al. (2021)).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>model</code>	object returned from the <code>rlmer</code> function
<code>args</code>	specification of function arguments
<code>result</code>	list with results, i.e., <code>call</code> for the the function call, <code>randeff</code> for the variance and correlation components, <code>coef</code> for the model coefficients, <code>weights</code> for the robustness weights, and <code>converge</code> for the convergence check, i.e., 1 = model converged, 0 = model singular, and -1 model not converged.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Koller, M. (2016). `robustlmm`: An R Package for Robust Estimation of Linear Mixed-Effects Models. *Journal of Statistical Software*, 75(6), 1–24. <https://doi.org/10.18637/jss.v075.i06>

See Also

[coeff.robust](#), [summa](#)

Examples

```
## Not run:
# Load lme4, lmerTest, and misty package
misty::libraries(lme4, lmerTest, misty)
```

```

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

#-----
# Multilevel and Linear Mixed-Effects Model

# Cluster-mean centering, center() from the misty package
Demo.twolevel <- center(Demo.twolevel, x2, type = "CWC", cluster = "cluster")

# Grand-mean centering, center() from the misty package
Demo.twolevel <- center(Demo.twolevel, w1, type = "CGM", cluster = "cluster")

# Estimate two-level mixed-effects model
mod.lmer2 <- lmer(y1 ~ x2.c + w1.c + x2.c:w1.c + (1 + x2.c | cluster), data = Demo.twolevel)

# Example 1a: Default setting
mod.lmer2r <- robust.lmer(mod.lmer2)

# Example 1b: Extract robustness weights
mod.lmer2r$result$weight$iresid
mod.lmer2r$result$weight$iraneff

#-----
# Write Results

# Example 2a: Write results into a text file
robust.lmer(mod.lmer2, write = "Robust_lmer.txt", output = FALSE)

# Example 2b: Write results into a Excel file
robust.lmer(mod.lmer2, write = "Robust_lmer.xlsx", output = FALSE)

## End(Not run)

```

script.copy

Save Copy of the Current Script in RStudio

Description

This function saves a copy of the current script in RStudio. By default, a folder called `_R_Script_Archive` will be created to save the copy of the current R script with the current date and time into the folder. Note that the current R script needs to have a file location before the script can be copied.

Usage

```

script.copy(file = NULL, folder = "_R_Script_Archive", create.folder = TRUE,
            time = TRUE, format = "%Y-%m-%d_%H%M", overwrite = TRUE,
            check = TRUE)

```

Arguments

file	a character string naming the file of the copy without the file extension ".R". By default, the file of the copy has the same name as the original file.
folder	a character string naming the folder in which the file of the copy is saved. If NULL, the file of the copy is saved in the same folder as the original file. By default, the file of the copy is saved into a folder called "_R_Script_Archive".
create.folder	logical: if TRUE (default), folder(s) specified in the file argument is created. If FALSE and the folder does not exist, then a error message is printed on the console.
time	logical: if TRUE (default), the current time is attached to the name of the file specified in the argument file.
format	a character string indicating the format if the POSIXct class resulting from the Sys.time function. The default setting provides a character string indicating the year, month, day, minutes, and seconds. See the help page of the format.POSIXct function.
overwrite	logical: if TRUE (default) an existing destination file is overwritten.
check	logical: if TRUE (default), argument specification is checked.

Note

This function uses the getSourceEditorContext() function in the **rstudioapi** package by Kevin Ushey, JJ Allaire, Hadley Wickham, and Gary Ritchie (2023).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Ushey, K., Allaire, J., Wickham, H., & Ritchie, G. (2023). *rstudioapi: Safely access the RStudio API*. R package version 0.15.0 <https://CRAN.R-project.org/package=rstudioapi>

See Also

[script.new](#), [script.close](#), [script.open](#), [script.save](#), [setsource](#)

Examples

```
## Not run:

# Example 1: Save copy current R script into the folder '_R_Script_Archive'
script.copy()

# Exmample 2: Save current R script as 'R_Script.R' into the folder 'Archive'
script.copy("R_Script", folder = "Archive", time = FALSE)

## End(Not run)
```

script.new*Open new R Script, R Markdown script, or SQL Script in RStudio*

Description

This function opens a new R script, R markdown script, or SQL script in RStudio.

Usage

```
script.new(text = "", type = c("r", "rmarkdown", "sql"),  
           position = rstudioapi::document_position(0, 0),  
           run = FALSE, check = TRUE)
```

Arguments

text	a character vector indicating what text should be inserted in the new R script. By default, an empty script is opened.
type	a character string indicating the type of document to be created, i.e., r (default) for an R script, rmarkdown for an R Markdown file, or sql for an SQL script.
position	document_position() function in the rstudioapi package indicating the cursor position.
run	logical: if TRUE, the code is executed after the document is created.
check	logical: if TRUE (default), argument specification is checked.

Note

This function uses the documentNew() function in the **rstudioapi** package by Kevin Ushey, JJ Allaire, Hadley Wickham, and Gary Ritchie (2023).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Ushey, K., Allaire, J., Wickham, H., & Ritchie, G. (2023). *rstudioapi: Safely access the RStudio API*. R package version 0.15.0 <https://CRAN.R-project.org/package=rstudioapi>

See Also

[script.close](#), [script.open](#), [script.save](#), [script.copy](#), [setsource](#)

Examples

```
## Not run:

# Example 1: Open new R script file
script.new()

# Example 2: Open new R script file and run some code
script.new("#-----")
# Example

# Generate 100 random numbers
rnorm(100)")

## End(Not run)
```

script.open

Open, Close and Save R Script in RStudio

Description

The function `script.open` opens an R script, R markdown script, or SQL script in RStudio, the function `script.close` closes an R script, and the function `script.save` saves an R script. Note that the R script need to have a file location before the script can be saved.

Usage

```
script.open(path, line = 1, col = 1, cursor = TRUE, run = FALSE,
            echo = TRUE, max.length = 999, spaced = TRUE, check = TRUE)

script.close(save = FALSE, check = TRUE)

script.save(all = FALSE, check = TRUE)
```

Arguments

<code>path</code>	a character string indicating the path of the script.
<code>line</code>	a numeric value indicating the line in the script to navigate to.
<code>col</code>	a numeric value indicating the column in the script to navigate to.
<code>cursor</code>	logical: if TRUE (default), the cursor moves to the requested location after opening the document.
<code>run</code>	logical: if TRUE, the code is executed after the document is opened
<code>echo</code>	logical: if TRUE (default), each expression is printed after parsing, before evaluation.
<code>max.length</code>	a numeric value indicating the maximal number of characters output for the deparse of a single expression.
<code>spaced</code>	logical: if TRUE (default), empty line is printed before each expression.

save	logical: if TRUE, the script is saved before closing when using the function <code>script.close</code> .
all	logical: if TRUE, all scripts opened in RStudio are saved when using the function <code>script.save</code> .
check	logical: if TRUE (default), argument specification is checked.

Note

This function uses the `documentOpen()`, `documentPath()`, `documentClose()`, `documentSave()`, and `documentSaveAll()` functions in the **rstudioapi** package by Kevin Ushey, JJ Allaire, Hadley Wickham, and Gary Ritchie (2023).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Ushey, K., Allaire, J., Wickham, H., & Ritchie, G. (2023). *rstudioapi: Safely access the RStudio API*. R package version 0.15.0 <https://CRAN.R-project.org/package=rstudioapi>

See Also

[script.save](#), [script.copy](#), [setsource](#)

Examples

```
## Not run:

# Example 1: Open R script file
script.open("script.R")

# Example 2: Open R script file and run the code
script.open("script.R", run = TRUE)

# Example 3: Close current R script file
script.close()

# Example 4: Save current R script
script.save()

# Example 5: Save all R scripts
script.save(all = TRUE)

## End(Not run)
```

setsource

Set Working Directory to the Source File Location

Description

This function sets the working directory to the source file location (i.e., path of the current R script) in RStudio and is equivalent to using the menu item Session - Set Working Directory - To Source File Location. Note that the R script needs to have a file location before this function can be used.

Usage

```
setsource(path = TRUE, check = TRUE)
```

Arguments

path	logical: if TRUE (default), the path of the source file is shown on the console.
check	logical: if TRUE, argument specification is checked.

Value

Returns the path of the source file location.

Note

This function uses the `documentPath()` function in the **rstudioapi** package by Kevin Ushey, JJ Allaire, Hadley Wickham, and Gary Ritchie (2023).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Ushey, K., Allaire, J., Wickham, H., & Ritchie, G. (2023). *rstudioapi: Safely access the RStudio API*. R package version 0.15.0 <https://CRAN.R-project.org/package=rstudioapi>

See Also

[script.close](#), [script.new](#), [script.open](#), [script.save](#)

Examples

```
## Not run:

# Example 1: Set working directory to the source file location
setsource()

# Example 2: Set working directory to the source file location
# and assign path to an object
path <- setsource()
path

## End(Not run)
```

size.mean	<i>Sample Size Determination</i>
-----------	----------------------------------

Description

This function performs sample size determination the one-sample and two-sample t-tests, proportions, and Pearson product-moment correlation coefficients based on precision requirements (i.e., type-I-risk, type-II-risk and an effect size).

Usage

```
size.mean(delta, sample = c("two.sample", "one.sample"),
          alternative = c("two.sided", "less", "greater"),
          alpha = 0.05, beta = 0.1, write = NULL, append = TRUE,
          check = TRUE, output = TRUE)

size.prop(pi = 0.5, delta, sample = c("two.sample", "one.sample"),
          alternative = c("two.sided", "less", "greater"),
          alpha = 0.05, beta = 0.1, correct = FALSE, write = NULL,
          append = TRUE, check = TRUE, output = TRUE)

size.cor(rho, delta,
         alternative = c("two.sided", "less", "greater"),
         alpha = 0.05, beta = 0.1, write = NULL, append = TRUE,
         check = TRUE, output = TRUE)
```

Arguments

delta	a numeric value indicating the minimum mean difference to be detected, δ .
sample	a character string specified in the function size.mean or size.prop specifying a one- or two-sample t-test or a proportion test, i.e., "two.sample" (default) for a two-sample test and "one.sample" for a one-sample test.
alternative	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".

alpha	a numeric value indicating the type-I-risk, α .
beta	a numeric value indicating the type-II-risk, β .
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown.
pi	a numeric value specified in the function size.prop indicating the true value of the probability under the null hypothesis in the one-sample test π_0 or a number indicating the true value of the probability in group 1 in the two-sample test π_1 .
rho	a numeric value specified in the function size.cor indicating the correlation coefficient under the null hypothesis ρ_0 .
correct	logical: if TRUE, continuity correction is applied.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>.

References

Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.

Rasch, D., Pilz, J., Verdooren, L. R., & Gebhardt, G. (2011). *Optimal experimental design with R*. Chapman & Hall/CRC.

See Also

[test.t](#), [prop.test](#), [cor.test](#), [cor.matrix](#)

Examples

```
#-----
# Example 1: One- and two-sample t-test

# Example 1a: One-sample t-test
# H0: mu = mu.0, H1: mu != mu.0
# alpha = 0.05, beta = 0.2, delta = 0.5
size.mean(delta = 0.5, sample = "one.sample",
           alternative = "two.sided", alpha = 0.05, beta = 0.2)

# Example 1b: One-sided two-sample test
# H0: mu.1 >= mu.2, H1: mu.1 < mu.2
# alpha = 0.01, beta = 0.1, delta = 1
size.mean(delta = 1, sample = "two.sample",
           alternative = "less", alpha = 0.01, beta = 0.1)

#-----
```

```
# Example 2: One- and two-sample test for proportions

# Example 2a: Two-sided one-sample test
# H0: pi = 0.5, H1: pi != 0.5
# alpha = 0.05, beta = 0.2, delta = 0.2
size.prop(pi = 0.5, delta = 0.2, sample = "one.sample",
          alternative = "two.sided", alpha = 0.05, beta = 0.2)

# Example 2b: One-sided two-sample test
# H0: pi.1 <= pi.1 = 0.5, H1: pi.1 > pi.2
# alpha = 0.01, beta = 0.1, delta = 0.2
size.prop(pi = 0.5, delta = 0.2, sample = "two.sample",
          alternative = "greater", alpha = 0.01, beta = 0.1)

#-----
# Example 3: Testing the Pearson product-moment correlation coefficient

# H0: rho = 0.3, H1: rho != 0.3
# alpha = 0.05, beta = 0.2, delta = 0.2
size.cor(rho = 0.3, delta = 0.2, alpha = 0.05, beta = 0.2)

# H0: rho <= 0.3, H1: rho > 0.3
# alpha = 0.05, beta = 0.2, delta = 0.2
size.cor(rho = 0.3, delta = 0.2,
          alternative = "greater", alpha = 0.05, beta = 0.2)
```

skewness

Univariate and Multivariate Skewness and Kurtosis

Description

The function `skewness` computes the univariate sample or population skewness and conduct's Mardia's test for multivariate skewness, while the function `kurtosis` computes the univariate sample or population (excess) kurtosis or the multivariate (excess) kurtosis and conduct's Mardia's test for multivariate kurtosis. By default, the function computes the sample univariate skewness or multivariate skewness and the univariate sample excess kurtosis or multivariate excess kurtosis.

Usage

```
skewness(data, ..., sample = TRUE, digits = 2, p.digits,
         as.na = NULL, check = TRUE, output = TRUE)

kurtosis(data, ..., sample = TRUE, center = TRUE, digits = 2, p.digits,
         as.na = NULL, check = TRUE, output = TRUE)
```

Arguments

`data` a numeric vector or data frame.

...	an expression indicating the variable names in data, e.g., <code>skewness(dat, x1)</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the <code>df.subset</code> function.
sample	logical: if TRUE (default), the univariate sample skewness or kurtosis is computed, while the population skewness or kurtosis is computed when <code>sample = FALSE</code> .
center	logical: if TRUE (default), the univariate or multivariate kurtosis is centered, so that the expected kurtosis under univariate or multivariate normality is 0, while the expected kurtosis under univariate or multivariate normality is 3 when <code>center = FALSE</code> .
digits	an integer value indicating the number of decimal places to be used. Note that this argument only applied when computing multivariate skewness and kurtosis.
p.digits	an integer value indicating the number of decimal places to be used for displaying the <i>p</i> -values.
as.na	a numeric vector indicating user-defined missing values, i.e., these values are converted to NA before conducting the analysis.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console. Note that this argument only applied when computing multivariate skewness and kurtosis.

Details

Univariate Skewness and Kurtosis Univariate skewness and kurtosis are computed based on the same formula as in SAS and SPSS:

- *Population Skewness*

$$\sqrt{n} \frac{\sum_{i=1}^n (X_i - \bar{X})^3}{(\sum_{i=1}^n (X_i - \bar{X})^2)^{3/2}}$$

- *Sample Skewness*

$$\frac{n\sqrt{n-1}}{n-2} \frac{\sum_{i=1}^n (X_i - \bar{X})^3}{(\sum_{i=1}^n (X_i - \bar{X})^2)^{3/2}}$$

- *Population Excess Kurtosis*

$$n \frac{\sum_{i=1}^n (X_i - \bar{X})^4}{(\sum_{i=1}^n (X_i - \bar{X})^2)^2} - 3$$

- *Sample Excess Kurtosis*

$$(n+1) \frac{\sum_{i=1}^n (X_i - \bar{X})^4}{(\sum_{i=1}^n (X_i - \bar{X})^2)^2} - 3 + 6 \frac{n-1}{(n-2)(n-3)}$$

Note that missing values (NA) are stripped before the computation and that at least 3 observations are needed to compute skewness and at least 4 observations are needed to compute kurtosis.

Multivariate Skewness and Kurtosis Mardia's multivariate skewness and kurtosis compares the joint distribution of several variables against a multivariate normal distribution. The expected skewness is 0 for a multivariate normal distribution, while the expected kurtosis is $p(p+2)$ for

a multivariate distribution of p variables. However, this function scales the multivariate kurtosis on $p(p + 2)$ according to the default setting `center = TRUE` so that the expected kurtosis under multivariate normality is 0. Multivariate skewness and kurtosis are tested for statistical significance based on the chi-square distribution for skewness and standard normal distribution for the kurtosis. If at least one of the tests is statistically significant, the underlying joint population is inferred to be non-normal. Note that non-significance of these statistical tests do not imply multivariate normality.

Value

Returns univariate skewness or kurtosis of data or an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>data</code>	a numeric vector or data frame specified in <code>data</code>
<code>args</code>	specification of function arguments
<code>result</code>	result table

Note

These functions implemented a modified copy of the `mardia()` function in the **psych** package by William Revelle (2024).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Cain, M. K., Zhang, Z., & Yuan, KH. (2024). Univariate and multivariate skewness and kurtosis for measuring nonnormality: Prevalence, influence and estimation. *Behavior Research Methods*, 49, 1716–1735. <https://doi.org/10.3758/s13428-016-0814-1>
- Mardia, K. V. (1970). Measures of multivariate skewness and kurtosis with applications. *Biometrika*, 57(3), 519-530. <https://doi.org/10.2307/2334770>
- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.
- William Revelle (2024). *psych: Procedures for Psychological, Psychometric, and Personality Research*. Northwestern University, Evanston, Illinois. R package version 2.4.6, <https://CRAN.R-project.org/package=psych>.

See Also

[descript](#)

Examples

```
# Example 1a: Compute univariate sample skewness
skewness(mtcars, mpg)

# Example 1b: Compute univariate sample excess kurtosis
kurtosis(mtcars, mpg)

# Example 2a: Compute multivariate skewness
skewness(mtcars)

# Example 2b: Compute multivariate excess kurtosis
kurtosis(mtcars)
```

summa	<i>Print Summary Output</i>
-------	-----------------------------

Description

This function prints a summary of the result object returned by the function "lm" for estimating linear regression models and for the result object returned by the function "lmer" from the **lme4** or **lmerTest** package, or by the function "rlmer" from the **robustlmm** package to estimate two- or three-level (robust) multilevel and linear mixed-effects models. By default, the function prints the function call, model summary, and the regression coefficient table.

Usage

```
summa(model, print = c("all", "default", "call", "descript", "cormat", "modsum",
                        "coef", "confint", "stdcoef", "vif"),
      robust = FALSE, ddf = c("Satterthwaite", "Kenward-Roger", "lme4"),
      conf.level = 0.95, method = c("profile", "wald", "boot"), R = 1000,
      boot = c("perc", "basic", "norm"), seed = NULL, digits = 2, p.digits = 3,
      write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

model	a fitted model of class "lm", "lmerMod", "rlmerMod", or "lmerModLmerTest"
print	a character vector indicating which results to print, i.e. "all", for all results, "call" for the function call, "descript" for descriptive statistics, cormat for the Pearson product-moment correlation matrix for models estimated by "lm" (see cor.matrix function) or within- and between-group correlation matrix for models estimated by "lmer" (see multilevel.cor function), modsum for the multiple correlation, r-squared, and F-test for models estimated by "lm" or model summary, marginal, and conditional R-squared for models estimated by "lmer" (see multilevel.r2 function), coef for the unstandardized coefficients for models estimated by "lm" or random effects and fixed effects for models estimated by "lmer", confint for the confidence interval for unstandardized coefficients, stdcoef for the standardized coefficients (see coeff.std function),

	and <code>vif</code> for the variance inflation factor (see <code>check.collin</code> function). The default setting is <code>print = c("call", "modsum", "coef")</code> . Note that when a fitted model of class <code>"r1merMod"</code> is specified for the argument <code>model</code> , the argument <code>print</code> is always <code>c("call", "coef")</code> , i.e., <code>"descript"</code> , <code>"cormat"</code> , <code>"modsum"</code> , <code>"confint"</code> , <code>"stdcoef"</code> , and <code>"vif"</code> are not available for an <code>"r1merMod"</code> object.
<code>robust</code>	logical: if TRUE, heteroscedasticity-consistent standard errors, confidence intervals, and heteroscedasticity-robust F-test using the HC4 estimator are computed for linear models estimated by using the <code>lm()</code> function or cluster-robust standard errors using the CR2 estimator is computed for multilevel and linear mixed-effects models estimated by using the <code>lmer()</code> function (see <code>coeff.robust</code> function).
<code>ddf</code>	a character string for specifying the method for computing the degrees of freedom when using the lmerTest package to obtain p -values for fixed effects in multilevel and linear mixed-effects models, i.e., <code>"Satterthwaite"</code> (default) for Satterthwaite's method, <code>"Kenward-Roger"</code> for the Kenward-Roger's method, and <code>"lme4"</code> for the lme4-summary without degrees of freedom and significance values (see Kuznetsova et al., 2017). Note that when a fitted model of class <code>"r1merMod"</code> is specified for the argument <code>model</code> , Satterthwaite or Kenward-Roger degrees of freedom are computed only if the R package lmerTest is loaded.
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.
<code>method</code>	a character string for specifying the method for computing confidence intervals (CI), i.e., <code>"profile"</code> (default) for computing a likelihood profile and finding the appropriate cutoffs based on the likelihood ratio test, <code>"Wald"</code> for approximating the CIs based on the estimated local curvature of the likelihood surface, and <code>"boot"</code> for performing bootstrapping with CIs computed from the bootstrap distribution according to the argument <code>boot</code> .
<code>R</code>	a numeric value indicating the number of bootstrap replicates (default is 1000).
<code>boot</code>	a character string for specifying the type of bootstrap confidence intervals (CI), i.e., <code>"perc"</code> (default), for the percentile bootstrap CI, <code>"basic"</code> for the basic bootstrap CI, and <code>"norm"</code> for the normal approximation bootstrap CI.
<code>seed</code>	a numeric value specifying seeds of the pseudo-random numbers used in the bootstrap algorithm when conducting bootstrapping.
<code>digits</code>	an integer value indicating the number of decimal places to be used.
<code>p.digits</code>	an integer value indicating the number of decimal places to be used for displaying multiple R, R-squared and p -value.
<code>write</code>	a character string naming a file for writing the output into either a text file with file extension <code>".txt"</code> (e.g., <code>"Output.txt"</code>) or Excel file with file extension <code>".xlsx"</code> (e.g., <code>"Output.xlsx"</code>). If the file name does not contain any file extension, an Excel file will be written.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension <code>.txt</code> specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.

Details

Robust Estimation of Multilevel and Linear Mixed-Effects Models The function `rlmer` from the **robustlmm** package does not provide any degrees of freedom or significance values. This function re-estimates the model without using robust estimation to obtain the Satterthwaite or Kenward-Roger degrees of freedom depending on the argument `ddf` before computing significance values for the regression coefficients based on parameter estimates and standard error of the robust multilevel mixed-effects (see Sleegers et al. (2021)).

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>model</code>	model specified in <code>model</code>
<code>args</code>	specification of function arguments
<code>result</code>	list with results, i.e., <code>call</code> for the the function call, <code>descript</code> for descriptive statistics, <code>cormat</code> for the correlation matrix, <code>modsum</code> for the model summary, <code>randeff</code> for the variance and correlation components, <code>coef</code> for the model coefficients, <code>weights</code> for the robustness weights and <code>converg</code> for the convergence check, i.e., 1 = model converged, 0 = model singular, and -1 model not converged.

Author(s)

Takuya Yanagida

References

- Kuznetsova, A, Brockhoff, P. B., & Christensen, R. H. B. (2017). lmerTest Package: Tests in linear mixed effects models. *Journal of Statistical Software*, 82 13, 1-26. <https://doi.org/10.18637/jss.v082.i13>
- Sleegers, W. W. A., Proulx, T., & van Beest, I. (2021). Pupillometry and hindsight bias: Physiological arousal predicts compensatory behavior. *Social Psychological and Personality Science*, 12(7), 1146–1154. <https://doi.org/10.1177/1948550620966153>

See Also

[descript](#), [cor.matrix](#), [coeff.std](#), [coeff.robust](#), [check.collin](#)

Examples

```
#-----
# Linear Model

# Estimate linear model
mod.lm <- lm(mpg ~ cyl + disp, data = mtcars)

# Example 1a: Default setting
```

```

summa(mod.lm)

# Example 1b: Heteroscedasticity-consistent standard errors
summa(mod.lm, robust = TRUE)

# Example 1c: Print all available results
summa(mod.lm, print = "all")

# Example 1d: Print default results plus standardized coefficient
summa(mod.lm, print = c("default", "stdcoef"))

## Not run:
#-----
# Multilevel and Linear Mixed-Effects Model

# Load lme4 and misty package
misty::libraries(lme4, misty)

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

#-----
## Two-Level Data

# Cluster-mean centering, center() from the misty package
Demo.twolevel <- center(Demo.twolevel, x2, type = "CWC", cluster = "cluster")

# Grand-mean centering, center() from the misty package
Demo.twolevel <- center(Demo.twolevel, w1, type = "CGM", cluster = "cluster")

# Estimate two-level mixed-effects model
mod.lmer2 <- lmer(y1 ~ x2.c + w1.c + x2.c:w1.c + (1 + x2.c | cluster),
                 data = Demo.twolevel)

# Example 2a: Default setting
summa(mod.lmer2)

# Example 2b: Print all available results
summa(mod.lmer2, print = "all")

# Example 2c: Print default results plus standardized coefficient
summa(mod.lmer2, print = c("default", "stdcoef"))

# Load lmerTest package
library(lmerTest)

# Re-estimate two-level model using the lme4 and lmerTest package
mod.lmer2 <- lmer(y1 ~ x2.c + w1.c + x2.c:w1.c + (1 + x2.c | cluster), data = Demo.twolevel)

# Example 2d: Default setting, Satterthwaite's method
summa(mod.lmer2)

# Example 2e: Kenward-Roger's method

```

```

summa(mod.lmer2, ddf = "Kenward-Roger")

# Example 2f: Cluster-robust standard errors
summa(mod.lmer2, robust = TRUE)

#-----
## Robust Estimation using the R package robustlmm

# Estimate two-level mixed-effects model
mod.lmer2r <- robustlmm::rlmer(y1 ~ x2.c + w1.c + x2.c:w1.c + (1 + x2.c | cluster),
                             data = Demo.twolevel)

# Example 2f: Default setting
summa(mod.lmer2r)

#-----
## Three-Level Data

# Create arbitrary three-level data
Demo.threelevel <- data.frame(Demo.twolevel, cluster2 = Demo.twolevel$cluster,
                             cluster3 = rep(1:10, each = 250))

# Cluster-mean centering, center() from the misty package
Demo.threelevel <- center(Demo.threelevel, x1, type = "CWC", cluster = c("cluster3", "cluster2"))

# Cluster-mean centering, center() from the misty package
Demo.threelevel <- center(Demo.threelevel, w1, type = "CWC", cluster = c("cluster3", "cluster2"))

# Estimate three-level model using the lme4 package
mod.lmer3 <- lmer(y1 ~ x1.c + w1.c + (1 | cluster3/cluster2), data = Demo.threelevel)

# Example 3a: Default setting
summa(mod.lmer3)

# Example 3b: Print all available results
summa(mod.lmer3, print = "all")

#-----
## Robust Estimation using the R package robustlmm

# Estimate three-level model using the lme4 package
mod.lmer3r <- robustlmm::rlmer(y1 ~ x1.c + w1.c + (1 | cluster3/cluster2),
                             data = Demo.threelevel)

# Example 3c: Default setting
summa(mod.lmer3r)

#-----
# Write Results

# Example 4a: Write Results into a text file
summa(mod.lm, print = "all", write = "Linear_Model.txt")

```

```
# Example 4b: Write Results into a Excel file
summa(mod.lm, print = "all", write = "Linear_Model.xlsx")

## End(Not run)
```

test.levene

Levene's Test for Homogeneity of Variance

Description

This function performs Levene's test for homogeneity of variance across two or more independent groups including a plot showing violins and boxplots representing the distribution of the outcome variable for each group.

Usage

```
test.levene(formula, data, method = c("median", "mean"), conf.level = 0.95,
             hypo = TRUE, descript = TRUE, digits = 2, p.digits = 3, as.na = NULL,
             plot = FALSE, violin = TRUE, box = TRUE, jitter = FALSE,
             violin.alpha = 0.3, violin.trim = FALSE, box.alpha = 0.2,
             box.width = 0.2, jitter.size = 1.25, jitter.width = 0.05,
             jitter.height = 0, jitter.alpha = 0.2, gray = FALSE, start = 0.9,
             end = 0.4, color = NULL, xlab = NULL, ylab = NULL, ylim = NULL,
             ybreaks = ggplot2::waiver(), title = "", subtitle = "",
             filename = NULL, width = NA, height = NA,
             units = c("in", "cm", "mm", "px"), dpi = 600, write = NULL,
             append = TRUE, check = TRUE, output = TRUE)
```

Arguments

formula	a formula of the form $y \sim \text{group}$ where y is a numeric variable giving the data values and group a numeric variable, character variable or factor with two or more than two values or factor levels giving the corresponding groups.
data	a matrix or data frame containing the variables in the formula formula.
method	a character string specifying the method to compute the center of each group, i.e. method = "median" (default) to compute the Levene's test based on the median (aka Brown-Forsythe test) or method = "mean" to compute the Levene's test based on the arithmetic mean.
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
hypo	logical: if TRUE (default), null and alternative hypothesis are shown on the console.
descript	logical: if TRUE (default), descriptive statistics are shown on the console.
digits	an integer value indicating the number of decimal places to be used for displaying results.
p.digits	an integer value indicating the number of decimal places to be used for displaying the p -value.

as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
plot	logical: if TRUE, a plot showing violins with boxplots is drawn.
violin	logical: if TRUE (default), violins are drawn.
box	logical: if TRUE (default), boxplots are drawn.
jitter	logical: if TRUE (default), jittered data points are drawn.
violin.alpha	a numeric value between 0 and 1 for specifying the alpha argument in the <code>geom_violin</code> function for controlling the opacity of the violins.
violin.trim	logical: if TRUE, the tails of the violins to the range of the data is trimmed.
box.alpha	a numeric value between 0 and 1 for specifying the alpha argument in the <code>geom_boxplot</code> function for controlling the opacity of the boxplots.
box.width	a numeric value indicating the width of the boxplots.
jitter.size	a numeric value indicating the size aesthetic for the jittered data points.
jitter.width	a numeric value indicating the amount of horizontal jitter.
jitter.height	a numeric value indicating the amount of vertical jitter.
jitter.alpha	a numeric value between 0 and 1 for specifying the alpha argument in the <code>geom_jitter</code> function for controlling the opacity of the jittered data points.
gray	logical: if TRUE, the plot is drawn in gray scale.
start	a numeric value between 0 and 1, graphical parameter to specify the gray value at the low end of the palette.
end	a numeric value between 0 and 1, graphical parameter to specify the gray value at the high end of the palette.
color	a character vector, indicating the color of the violins and the boxes. By default, default ggplot2 colors are used.
xlab	a character string specifying the labels for the x-axis.
ylab	a character string specifying the labels for the y-axis.
ylim	a numeric vector of length two specifying limits of the limits of the y-axis.
ybreaks	a numeric vector specifying the points at which tick-marks are drawn at the y-axis.
title	a character string specifying the text for the title for the plot.
subtitle	a character string specifying the text for the subtitle for the plot.
filename	a character string indicating the filename argument (default is "NA_Pattern.pdf") including the file extension for the <code>ggsave</code> function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the file argument.
width	a numeric value indicating the width argument (default is the size of the current graphics device) for the <code>ggsave</code> function.
height	a numeric value indicating the height argument (default is the size of the current graphics device) for the <code>ggsave</code> function.
units	a character string indicating the units argument (default is in) for the <code>ggsave</code> function.

dpi	a numeric value indicating the dpi argument (default is 600) for the ggsave function.
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown.

Details

Levene's test is equivalent to a one-way analysis of variance (ANOVA) with the absolute deviations of observations from the mean of each group as dependent variable (center = "mean"). Brown and Forsythe (1974) modified the Levene's test by using the absolute deviations of observations from the median (center = "median"). By default, the Levene's test uses the absolute deviations of observations from the median.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
data	data frame specified in data
formula	formula of the current analysis
args	specification of function arguments
plot	ggplot2 object for plotting the results
result	list with result tables, i.e., <code>descript</code> for descriptive statistics and test for the ANOVA table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Brown, M. B., & Forsythe, A. B. (1974). Robust tests for the equality of variances. *Journal of the American Statistical Association*, 69, 364-367.

Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.

See Also

[aov.b](#), [test.t](#), [test.welch](#)

Examples

```
# Example 1: Levene's test based on the median
test.levene(mpg ~ gear, data = mtcars)

# Example 2: Levene's test based on the arithmetic mean
test.levene(mpg ~ gear, data = mtcars, method = "mean")

# Example 3: Levene's test based on the median, plot results
test.levene(mpg ~ gear, data = mtcars, plot = TRUE)

## Not run:
# Example 4: Write results into a text file
test.levene(mpg ~ gear, data = mtcars, write = "Levene.txt")

# Example 5: Levene's test based on the median, save plot
test.levene(mpg ~ gear, data = mtcars, plot = TRUE,
            filename = "Levene-test.png", dpi = 600, width = 6, height = 5)

## End(Not run)
```

test.t	<i>t-Test</i>
--------	---------------

Description

This function performs one-sample, two-sample, and paired-sample t-tests and provides descriptive statistics, effect size measure, and a plot showing bar plots with error bars for (difference-adjusted) confidence intervals.

Usage

```
test.t(x, ...)
```

```
## Default S3 method:
test.t(x, y = NULL, mu = 0, paired = FALSE,
       alternative = c("two.sided", "less", "greater"), conf.level = 0.95,
       hypo = TRUE, descript = TRUE, effsize = FALSE, weighted = FALSE,
       cor = TRUE, ref = NULL, correct = FALSE, digits = 2, p.digits = 3,
       as.na = NULL, plot = FALSE, bar = TRUE, point = FALSE, ci = TRUE,
       line = TRUE, jitter = FALSE, adjust = TRUE, point.size = 4, errorbar.width = 0.1,
       xlab = NULL, ylab = NULL, ylim = NULL, ybreaks = ggplot2::waiver(),
       linetype = 3, linewidth = 0.8, jitter.size = 1.25, jitter.width = 0.05,
       jitter.height = 0, jitter.alpha = 0.1, title = "",
       subtitle = "Confidence Interval", filename = NULL, width = NA, height = NA,
       units = c("in", "cm", "mm", "px"),
       dpi = 600, write = NULL, append = TRUE, check = TRUE, output = TRUE, ...)
```

```
## S3 method for class 'formula'
```

```
test.t(formula, data, alternative = c("two.sided", "less", "greater"),
  conf.level = 0.95, hypo = TRUE, descript = TRUE, effsize = FALSE, weighted = FALSE,
  cor = TRUE, ref = NULL, correct = FALSE, digits = 2, p.digits = 3, as.na = NULL,
  plot = FALSE, bar = TRUE, point = FALSE, ci = TRUE, line = TRUE,
  jitter = FALSE, adjust = TRUE, point.size = 4, errorbar.width = 0.1, xlab = NULL,
  ylab = NULL, ylim = NULL, ybreaks = ggplot2::waiver(), linetype = 3,
  linewidth = 0.8, jitter.size = 1.25, jitter.width = 0.05, jitter.height = 0,
  jitter.alpha = 0.1, title = "", subtitle = "Confidence Interval", filename = NULL,
  width = NA, height = NA, units = c("in", "cm", "mm", "px"), dpi = 600,
  write = NULL, append = TRUE, check = TRUE, output = TRUE, ...)
```

Arguments

<code>x</code>	a numeric vector of data values.
<code>...</code>	further arguments to be passed to or from methods.
<code>y</code>	a numeric vector of data values.
<code>mu</code>	a numeric value indicating the population mean under the null hypothesis. Note that the argument <code>mu</code> is only used when computing a one sample t-test.
<code>paired</code>	logical: if TRUE, paired-samples t-test is computed.
<code>alternative</code>	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
<code>hypo</code>	logical: if TRUE (default), null and alternative hypothesis are shown on the console.
<code>descript</code>	logical: if TRUE (default), descriptive statistics are shown on the console.
<code>effsize</code>	logical: if TRUE, effect size measure Cohen's d is shown on the console, see cohens.d function.
<code>weighted</code>	logical: if TRUE, the weighted pooled standard deviation is used to compute Cohen's d for a two-sample design (i.e., <code>paired = FALSE</code>), while standard deviation of the difference scores is used to compute Cohen's d for a paired-sample design (i.e., <code>paired = TRUE</code>).
<code>cor</code>	logical: if TRUE (default), <code>paired = TRUE</code> , and <code>weighted = FALSE</code> , Cohen's d for a paired-sample design while controlling for the correlation between the two sets of measurement is computed. Note that this argument is only used in a paired-sample design (i.e., <code>paired = TRUE</code>) when specifying <code>weighted = FALSE</code> .
<code>ref</code>	character string "x" or "y" for specifying the reference reference group when using the default <code>test.t()</code> function or a numeric value or character string indicating the reference group in a two-sample design when using the formula <code>test.t()</code> function. The standard deviation of the reference variable or reference group is used to standardized the mean difference to compute Cohen's d. Note that this argument is only used in a two-sample design (i.e., <code>paired = FALSE</code>).
<code>correct</code>	logical: if TRUE, correction factor to remove positive bias in small samples is used.
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.

<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying descriptive statistics and confidence interval.
<code>p.digits</code>	an integer value indicating the number of decimal places to be used for displaying the p -value.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>plot</code>	logical: if TRUE, a plot showing bar plots with error bars for confidence intervals is drawn.
<code>bar</code>	logical: if TRUE (default), bars representing means for each groups are drawn.
<code>point</code>	logical: if TRUE, points representing means for each groups are drawn.
<code>ci</code>	logical: if TRUE (default), error bars representing confidence intervals are drawn.
<code>jitter</code>	logical: if TRUE, jittered data points are drawn.
<code>line</code>	logical: if TRUE (default), a horizontal line is drawn at μ for the one-sample t-test or at 0 for the paired-sample t-test.
<code>adjust</code>	logical: if TRUE (default), difference-adjustment for the confidence intervals in a two-sample design is applied.
<code>point.size</code>	a numeric value indicating the size aesthetic for the point representing the mean value.
<code>errorbar.width</code>	a numeric value indicating the horizontal bar width of the error bar.
<code>xlab</code>	a character string specifying the labels for the x-axis.
<code>ylab</code>	a character string specifying the labels for the y-axis.
<code>ylim</code>	a numeric vector of length two specifying limits of the limits of the y-axis.
<code>ybreaks</code>	a numeric vector specifying the points at which tick-marks are drawn at the y-axis.
<code>linetype</code>	an integer value or character string specifying the line type for the line representing the population mean under the null hypothesis, i.e., 0 = blank, 1 = solid, 2 = dashed, 3 = dotted, 4 = dotdash, 5 = longdash, 6 = twodash.
<code>linewidth</code>	a numeric value indicating the linewidth aesthetic for the line representing the population mean under the null hypothesis.
<code>jitter.size</code>	a numeric value indicating the size aesthetic
<code>jitter.width</code>	a numeric value indicating the amount of horizontal jitter.
<code>jitter.height</code>	a numeric value indicating the amount of vertical jitter.
<code>jitter.alpha</code>	a numeric value between 0 and 1 for specifying the alpha argument in the <code>geom_jitter</code> function for controlling the opacity of the jittered data points.
<code>title</code>	a character string specifying the text for the title for the plot.
<code>subtitle</code>	a character string specifying the text for the subtitle for the plot.
<code>filename</code>	a character string indicating the filename argument (default is "NA_Pattern.pdf") including the file extension for the <code>ggsave</code> function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the <code>file</code> argument.

width	a numeric value indicating the width argument (default is the size of the current graphics device) for the ggsave function.
height	a numeric value indicating the height argument (default is the size of the current graphics device) for the ggsave function.
units	a character string indicating the units argument (default is in) for the ggsave function.
dpi	a numeric value indicating the dpi argument (default is 600) for the ggsave function.
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.
formula	in case of two sample t-test (i.e., paired = FALSE), a formula of the form $y \sim \text{group}$ where group is a numeric variable, character variable or factor with two values or factor levels giving the corresponding groups.
data	a matrix or data frame containing the variables in the formula formula.

Details

Effect Size Measure By default, Cohen's d based on the non-weighted standard deviation (i.e., `weighted = FALSE`) which does not assume homogeneity of variance is computed (see Delacre et al., 2021) when requesting an effect size measure (i.e., `effsize = TRUE`). Cohen's d based on the pooled standard deviation assuming equality of variances between groups can be requested by specifying `weighted = TRUE`.

Value

Returns an object of class `misty.object`, which is a list with following entries:

call	function call
type	type of analysis
sample	type of sample, i.e., one-, two-, or paired sample
formula	formula
data	data frame with the outcome and grouping variable
args	specification of function arguments
plot	ggplot2 object for plotting the results
result	result table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.
- Delacre, M., Lakens, D., Ley, C., Liu, L., & Leys, C. (2021). Why Hedges' g 's based on the non-pooled standard deviation should be reported with Welch's t -test. <https://doi.org/10.31234/osf.io/tu6mp>

See Also

[aov.b](#), [aov.w](#), [test.welch](#), [test.z](#), [test.levene](#), [cohens.d](#), [ci.mean.diff](#), [ci.mean](#)

Examples

```
#-----
# One-Sample Design

# Example 1a: Two-sided one-sample t-test, population mean = 20
test.t(mtcars$mpg, mu = 20)

# Example 1b: One-sided one-sample t-test, population mean = 20, print Cohen's d
test.t(mtcars$mpg, mu = 20, alternative = "greater", effsize = TRUE)

# Example 1c: Two-sided one-sample t-test, population mean = 20, plot results
test.t(mtcars$mpg, mu = 20, plot = TRUE)

## Not run:
# Example 1d: Two-sided one-sample t-test, population mean = 20, save plot
test.t(mtcars$mpg, mu = 20, plot = TRUE, filename = "One-sample_t-test.png",
      width = 4, height = 5)

## End(Not run)

#-----
# Two-Sample Design

# Example 2a: Two-sided two-sample t-test
test.t(mpg ~ vs, data = mtcars)

# Example 2b: Two-sided two-sample t-test, alternative specification
test.t(c(3, 1, 4, 2, 5, 3, 6, 7), c(5, 2, 4, 3, 1))

# Example 2c: One-sided two-sample t-test, print Cohen's d with weighted pooled SD
test.t(mpg ~ vs, data = mtcars, alternative = "greater", effsize = TRUE)

# Example 2d: Two-sided two-sample t-test, plot results
test.t(mpg ~ vs, data = mtcars, plot = TRUE)

## Not run:
# Example 2e: Two-sided two-sample t-test, plot results
test.t(mpg ~ vs, data = mtcars, plot = TRUE, filename = "Two-sample_t-test.png",
      width = 5, height = 6)
```

```
## End(Not run)

#-----
# Paired-Sample Design

# Example 3a: Two-sided paired-sample t-test
test.t(mtcars$drat, mtcars$wt, paired = TRUE)

# Example 3b: One-sided paired-sample t-test,
# print Cohen's d based on the SD of the difference scores
test.t(mtcars$drat, mtcars$wt, paired = TRUE, alternative = "greater",
       effsize = TRUE)

# Example 3c: Two-sided paired-sample t-test, plot results
test.t(mtcars$drat, mtcars$wt, paired = TRUE, plot = TRUE)

## Not run:
# Example 3d: Two-sided paired-sample t-test, save plot
test.t(mtcars$drat, mtcars$wt, paired = TRUE, plot = TRUE,
       filename = "Paired-sample_t-test.png", width = 4, height = 5)

## End(Not run)
```

test.welch

Welch's Test

Description

This function performs Welch's two-sample t-test and Welch's ANOVA including Games-Howell post hoc test for multiple comparison and provides descriptive statistics, effect size measures, and a plot showing bars representing means for each group and error bars for difference-adjusted confidence intervals.

Usage

```
test.welch(formula, data, alternative = c("two.sided", "less", "greater"),
           posthoc = FALSE, conf.level = 0.95, hypo = TRUE, descript = TRUE,
           effsize = FALSE, weighted = FALSE, ref = NULL, correct = FALSE,
           digits = 2, p.digits = 3, as.na = NULL, plot = FALSE, bar = TRUE,
           point = FALSE, ci = TRUE, jitter = FALSE, adjust = TRUE,
           point.size = 3, errorbar.width = 0.1, jitter.size = 1.25,
           jitter.width = 0.05, jitter.height = 0, jitter.alpha = 0.1,
           xlab = NULL, ylab = "y", ylim = NULL, ybreaks = ggplot2::waiver(),
           title = NULL, subtitle = "Confidence Interval", filename = NULL,
           width = NA, height = NA, units = c("in", "cm", "mm", "px"),
           dpi = 600, write = NULL, append = TRUE, check = TRUE, output = TRUE)
```

Arguments

formula	a formula of the form $y \sim \text{group}$ where y is a numeric variable giving the data values and group a numeric variable, character variable or factor with two or more than two values or factor levels giving the corresponding groups.
data	a matrix or data frame containing the variables in the formula formula.
alternative	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less". Note that this argument is only used when conducting Welch's two-sample t-test.
posthoc	logical: if TRUE, Games-Howell post hoc test for multiple comparison is conducted when performing Welch's ANOVA.
conf.level	a numeric value between 0 and 1 indicating the confidence level of the interval.
hypo	logical: if TRUE (default), null and alternative hypothesis are shown on the console.
descript	logical: if TRUE (default), descriptive statistics are shown on the console.
effsize	logical: if TRUE, effect size measure Cohen's d for Welch's two-sample t-test (see cohens.d), η^2 and ω^2 for Welch's ANOVA and Cohen's d for the post hoc tests are shown on the console.
weighted	logical: if TRUE, the weighted pooled standard deviation is used to compute Cohen's d .
ref	a numeric value or character string indicating the reference group. The standard deviation of the reference group is used to standardized the mean difference to compute Cohen's d .
correct	logical: if TRUE, correction factor to remove positive bias in small samples is used.
digits	an integer value indicating the number of decimal places to be used for displaying descriptive statistics and confidence interval.
p.digits	an integer value indicating the number of decimal places to be used for displaying the p -value.
as.na	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
plot	logical: if TRUE, a plot showing error bars for confidence intervals is drawn.
bar	logical: if TRUE (default), bars representing means for each groups are drawn.
point	logical: if TRUE, points representing means for each groups are drawn.
ci	logical: if TRUE (default), error bars representing confidence intervals are drawn.
jitter	logical: if TRUE, jittered data points are drawn.
adjust	logical: if TRUE (default), difference-adjustment for the confidence intervals is applied.
point.size	a numeric value indicating the size aesthetic for the point representing the mean value.
errorbar.width	a numeric value indicating the horizontal bar width of the error bar.
jitter.size	a numeric value indicating the size aesthetic for the jittered data points.

jitter.width	a numeric value indicating the amount of horizontal jitter.
jitter.height	a numeric value indicating the amount of vertical jitter.
jitter.alpha	a numeric value between 0 and 1 for specifying the alpha argument in the geom_jitter function for controlling the opacity of the jittered data points.
xlab	a character string specifying the labels for the x-axis.
ylab	a character string specifying the labels for the y-axis.
ylim	a numeric vector of length two specifying limits of the limits of the y-axis.
ybreaks	a numeric vector specifying the points at which tick-marks are drawn at the y-axis.
title	a character string specifying the text for the title of the plot.
subtitle	a character string specifying the text for the subtitle of the plot.
filename	a character string indicating the filename argument including the file extension in the ggsave function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the filename argument. Note that plots can only be saved when plot = TRUE.
width	a numeric value indicating the width argument (default is the size of the current graphics device) in the ggsave function.
height	a numeric value indicating the height argument (default is the size of the current graphics device) in the ggsave function.
units	a character string indicating the units argument (default is in) in the ggsave function.
dpi	a numeric value indicating the dpi argument (default is 600) in the ggsave function.
write	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

Effect Size Measure By default, Cohen's d based on the non-weighted standard deviation (i.e., weighted = FALSE) which does not assume homogeneity of variance is computed (see Delacre et al., 2021) when requesting an effect size measure (i.e., effsize = TRUE). Cohen's d based on the pooled standard deviation assuming equality of variances between groups can be requested by specifying weighted = TRUE.

Value

Returns an object of class misty.object, which is a list with following entries:

call	function call
------	---------------

type	type of analysis
sample	type of sample, i.e., one-, two-, or paired sample
data	data frame with the outcome and grouping variable
formula	formula
args	specification of function arguments
plot	ggplot2 object for plotting the results
result	result table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.
- Delacre, M., Lakens, D., Ley, C., Liu, L., & Leys, C. (2021). Why Hedges' g 's based on the non-pooled standard deviation should be reported with Welch's t -test. <https://doi.org/10.31234/osf.io/tu6mp>

See Also

[test.t](#), [test.z](#), [test.levene](#), [aov.b](#), [cohens.d](#), [ci.mean.diff](#), [ci.mean](#)

Examples

```
#-----
# Two-Sample Design

# Example 1a: Two-sided two-sample Welch-test
test.welch(mpg ~ vs, data = mtcars)

# Example 1b: One-sided two-sample Welch-test
test.welch(mpg ~ vs, data = mtcars, alternative = "greater")

# Example 1c: Two-sided two-sample Welch-test, print Cohen's d
test.welch(mpg ~ vs, data = mtcars, effsize = TRUE)

# Example 1d: Two-sided two-sample Welch-test, plot results
test.welch(mpg ~ vs, data = mtcars, plot = TRUE)

#-----
# Multiple-Sample Design

# Example 2a: Welch's ANOVA
test.welch(mpg ~ gear, data = mtcars)

# Example 2b: Welch's ANOVA, Games-Howell post hoc test
test.welch(mpg ~ gear, data = mtcars, posthoc = TRUE)
```

```
# Example 2c: Welch's ANOVA, print eta-squared and omega-squared
test.welch(mpg ~ gear, data = mtcars, effsize = TRUE)

# Example 2d: Welch's ANOVA, plot results
test.welch(mpg ~ gear, data = mtcars, plot = TRUE)

## Not run:
# Example 2e: Welch's ANOVA, save plot
test.welch(mpg ~ gear, data = mtcars, plot = TRUE,
           filename = "Multiple-sample_Welch-test.png", width = 6, height = 5)

## End(Not run)
```

test.z	<i>z-Test</i>
--------	---------------

Description

This function performs one-sample, two-sample, and paired-sample z-tests and provides descriptive statistics, effect size measure, and a plot showing error bars for (difference-adjusted) confidence intervals with jittered data points.

Usage

```
test.z(x, ...)
```

Default S3 method:

```
test.z(x, y = NULL, sigma = NULL, sigma2 = NULL, mu = 0,
       paired = FALSE, alternative = c("two.sided", "less", "greater"),
       conf.level = 0.95, hypo = TRUE, descript = TRUE, effsize = FALSE,
       digits = 2, p.digits = 3, as.na = NULL, plot = FALSE, bar = TRUE,
       point = FALSE, ci = TRUE, line = TRUE, jitter = FALSE, adjust = TRUE,
       point.size = 4, errorbar.width = 0.1, xlab = NULL, ylab = NULL,
       ylim = NULL, ybreaks = ggplot2::waiver(), linetype = 3, linewidth = 0.8,
       jitter.size = 1.25, jitter.width = 0.05, jitter.height = 0,
       jitter.alpha = 0.1, title = "", subtitle = "Confidence Interval",
       filename = NULL, width = NA, height = NA, units = c("in", "cm", "mm", "px"),
       dpi = 600, write = NULL, append = TRUE, check = TRUE,
       output = TRUE, ...)
```

S3 method for class 'formula'

```
test.z(formula, data, sigma = NULL, sigma2 = NULL,
       alternative = c("two.sided", "less", "greater"), conf.level = 0.95,
       hypo = TRUE, descript = TRUE, effsize = FALSE, digits = 2, p.digits = 3,
       as.na = NULL, plot = FALSE, bar = TRUE, point = FALSE, ci = TRUE,
       line = TRUE, jitter = FALSE, adjust = TRUE, point.size = 4, errorbar.width = 0.1,
       xlab = NULL, ylab = NULL, ylim = NULL, ybreaks = ggplot2::waiver(),
       linetype = 3, linewidth = 0.8, jitter.size = 1.25, jitter.width = 0.05,
```

```
jitter.height = 0, jitter.alpha = 0.1, title = "",
subtitle = "Confidence Interval", filename = NULL, width = NA, height = NA,
units = c("in", "cm", "mm", "px"), dpi = 600, write = NULL, append = TRUE,
check = TRUE, output = TRUE, ...)
```

Arguments

<code>x</code>	a numeric vector of data values.
<code>...</code>	further arguments to be passed to or from methods.
<code>y</code>	a numeric vector of data values.
<code>sigma</code>	a numeric vector indicating the population standard deviation(s). In case of two-sample z-test, equal standard deviations are assumed when specifying one value for the argument <code>sigma</code> ; when specifying two values for the argument <code>sigma</code> , unequal standard deviations are assumed. Note that either argument <code>sigma</code> or argument <code>sigma2</code> is specified.
<code>sigma2</code>	a numeric vector indicating the population variance(s). In case of two-sample z-test, equal variances are assumed when specifying one value for the argument <code>sigma2</code> ; when specifying two values for the argument <code>sigma</code> , unequal variance are assumed. Note that either argument <code>sigma</code> or argument <code>sigma2</code> is specified.
<code>mu</code>	a numeric value indicating the population mean under the null hypothesis. Note that the argument <code>mu</code> is only used when computing a one-sample z-test.
<code>paired</code>	logical: if TRUE, paired-sample z-test is computed.
<code>alternative</code>	a character string specifying the alternative hypothesis, must be one of "two.sided" (default), "greater" or "less".
<code>hypo</code>	logical: if TRUE (default), null and alternative hypothesis are shown on the console.
<code>descript</code>	logical: if TRUE (default), descriptive statistics are shown on the console.
<code>effsize</code>	logical: if TRUE, effect size measure Cohen's d is shown on the console.
<code>conf.level</code>	a numeric value between 0 and 1 indicating the confidence level of the interval.
<code>digits</code>	an integer value indicating the number of decimal places to be used for displaying descriptive statistics and confidence interval.
<code>p.digits</code>	an integer value indicating the number of decimal places to be used for displaying the <i>p</i> -value.
<code>as.na</code>	a numeric vector indicating user-defined missing values, i.e. these values are converted to NA before conducting the analysis.
<code>plot</code>	logical: if TRUE, a plot showing bar plots with error bars for confidence intervals is drawn.
<code>bar</code>	logical: if TRUE (default), bars representing means for each groups are drawn.
<code>point</code>	logical: if TRUE, points representing means for each groups are drawn.
<code>ci</code>	logical: if TRUE (default), error bars representing confidence intervals are drawn.
<code>jitter</code>	logical: if TRUE, jittered data points are drawn.
<code>line</code>	logical: if TRUE (default), a horizontal line is drawn at <code>mu</code> for the one-sample z-test or at 0 for the paired-sample z-test.

<code>adjust</code>	logical: if TRUE (default), difference-adjustment for the confidence intervals in a two-sample design is applied.
<code>point.size</code>	a numeric value indicating the size aesthetic for the point representing the mean value.
<code>errorbar.width</code>	a numeric value indicating the horizontal bar width of the error bar.
<code>xlab</code>	a character string specifying the labels for the x-axis.
<code>ylab</code>	a character string specifying the labels for the y-axis.
<code>ylim</code>	a numeric vector of length two specifying limits of the limits of the y-axis.
<code>ybreaks</code>	a numeric vector specifying the points at which tick-marks are drawn at the y-axis.
<code>linetype</code>	an integer value or character string specifying the line type for the line representing the population mean under the null hypothesis, i.e., 0 = blank, 1 = solid, 2 = dashed, 3 = dotted, 4 = dotdash, 5 = longdash, 6 = twodash.
<code>linewidth</code>	a numeric value indicating the linewidth aesthetic for the line representing the population mean under the null hypothesis.
<code>jitter.size</code>	a numeric value indicating the size aesthetic
<code>jitter.width</code>	a numeric value indicating the amount of horizontal jitter.
<code>jitter.height</code>	a numeric value indicating the amount of vertical jitter.
<code>jitter.alpha</code>	a numeric value between 0 and 1 for specifying the alpha argument in the <code>geom_jitter</code> function for controlling the opacity of the jittered data points.
<code>title</code>	a character string specifying the text for the title for the plot.
<code>subtitle</code>	a character string specifying the text for the subtitle for the plot.
<code>filename</code>	a character string indicating the filename argument (default is "NA_Pattern.pdf") including the file extension for the <code>ggsave</code> function. Note that one of ".eps", ".ps", ".tex", ".pdf" (default), ".jpeg", ".tiff", ".png", ".bmp", ".svg" or ".wmf" needs to be specified as file extension in the file argument.
<code>width</code>	a numeric value indicating the width argument (default is the size of the current graphics device) for the <code>ggsave</code> function.
<code>height</code>	a numeric value indicating the height argument (default is the size of the current graphics device) for the <code>ggsave</code> function.
<code>units</code>	a character string indicating the units argument (default is in) for the <code>ggsave</code> function.
<code>dpi</code>	a numeric value indicating the dpi argument (default is 600) for the <code>ggsave</code> function.
<code>write</code>	a character string naming a text file with file extension ".txt" (e.g., "Output.txt") for writing the output into a text file.
<code>append</code>	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in <code>write</code> , if FALSE existing text file will be overwritten.
<code>check</code>	logical: if TRUE (default), argument specification is checked.
<code>output</code>	logical: if TRUE (default), output is shown on the console.
<code>formula</code>	in case of two sample z-test (i.e., <code>paired = FALSE</code>), a formula of the form $y \sim \text{group}$ where <code>group</code> is a numeric variable, character variable or factor with two values or factor levels giving the corresponding groups.
<code>data</code>	a matrix or data frame containing the variables in the formula <code>formula</code> .

Details

Cohen's d reported when argument `effsize = TRUE` is based on the population standard deviation specified in `sigma` or the square root of the population variance specified in `sigma2`. In a one-sample and paired-sample design, Cohen's d is the mean of the difference scores divided by the population standard deviation of the difference scores (i.e., equivalent to Cohen's d_z according to Lakens, 2013). In a two-sample design, Cohen's d is the difference between means of the two groups of observations divided by either the population standard deviation when assuming and specifying equal standard deviations or the unweighted pooled population standard deviation when assuming and specifying unequal standard deviations.

Value

Returns an object of class `misty.object`, which is a list with following entries:

<code>call</code>	function call
<code>type</code>	type of analysis
<code>sample</code>	type of sample, i.e., one-, two-, or paired sample
<code>formula</code>	formula
<code>data</code>	data frame with the outcome and grouping variable
<code>args</code>	specification of function arguments
<code>plot</code>	ggplot2 object for plotting the results
<code>result</code>	result table

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Lakens, D. (2013). Calculating and reporting effect sizes to facilitate cumulative science: A practical primer for t-tests and ANOVAs. *Frontiers in Psychology*, 4, 1-12. <https://doi.org/10.3389/fpsyg.2013.00863>

Rasch, D., Kubinger, K. D., & Yanagida, T. (2011). *Statistics in psychology - Using R and SPSS*. John Wiley & Sons.

See Also

[test.t](#), [aov.b](#), [aov.w](#), [test.welch](#), [cohens.d](#), [ci.mean.diff](#), [ci.mean](#)

Examples

```
#-----
# One-Sample Design

# Example 1a: Two-sided one-sample z-test, population mean = 20, population SD = 6
test.z(mtcars$mpg, sigma = 6, mu = 20)

# Example 1b: One-sided one-sample z-test, population mean = 20, population SD = 6,
```

```

# print Cohen's d
test.z(mtcars$mpg, sigma = 6, mu = 20, alternative = "greater", effsize = TRUE)

# Example 1c: Two-sided one-sample z-test, population mean = 20, population SD = 6,
# plot results
test.z(mtcars$mpg, sigma = 6, mu = 20, plot = TRUE)

## Not run:
# Example 1d: Two-sided one-sample z-test, save plot
test.z(mtcars$mpg, sigma = 6, mu = 20, plot = TRUE, filename = "One-sample_z-test.png",
       width = 4, height = 5)

## End(Not run)

#-----
# Two-Sample Design

# Example 2a: Two-sided two-sample z-test, population SD = 6, equal SD assumption
test.z(mpg ~ vs, data = mtcars, sigma = 6)

# Example 2b: Two-sided two-sample z-test, alternative specification
test.z(c(3, 1, 4, 2, 5, 3, 6, 7), c(5, 2, 4, 3, 1), sigma = 1.2)

# Example 2c: Two-sided two-sample z-test, population SD = 4 and 6, unequal SD assumption
test.z(mpg ~ vs, data = mtcars, sigma = c(4, 6))

# Example 2d: One-sided two-sample z-test, population SD = 4 and 6, unequal SD assumption
# print Cohen's d
test.z(mpg ~ vs, data = mtcars, sigma = c(4, 6), alternative = "greater",
       effsize = TRUE)

# Example 2e: Two-sided two-sample z-test, population SD = 6, equal SD assumption
# plot results
test.z(mpg ~ vs, data = mtcars, sigma = 6, plot = TRUE)

## Not run:
# Example 2f: Two-sided two-sample z-test, save plot
test.z(mpg ~ vs, data = mtcars, sigma = 6, plot = TRUE, filename = "Two-sample_z-test.png",
       width = 5, height = 6)

## End(Not run)

#-----
# Paired-Sample Design

# Example 3a: Two-sided paired-sample z-test, population SD of difference score = 1.2
test.z(mtcars$drat, mtcars$wt, sigma = 1.2, paired = TRUE)

# Example 3b: One-sided paired-sample z-test, population SD of difference score = 1.2,
# print Cohen's d
test.z(mtcars$drat, mtcars$wt, sigma = 1.2, paired = TRUE,
       alternative = "greater", effsize = TRUE)

```

```
# Example 3c: Two-sided paired-sample z-test, population SD of difference score = 1.2,
# plot results
test.z(mtcars$drat, mtcars$wt, sigma = 1.2, paired = TRUE, plot = TRUE)

## Not run:
# Example 3d: Two-sided paired-sample z-test, save plot
test.z(mtcars$drat, mtcars$wt, sigma = 1.2, paired = TRUE, plot = TRUE,
       filename = "Paired-sample_z-test.png", width = 4, height = 5)

## End(Not run)
```

uniq

Extract Unique Elements and Count Number of Unique Elements

Description

The function `uniq` returns a vector, list or data frame with duplicated elements removed. By default, the function prints a data frame with missing values omitted and unique elements sorted increasing. The function `uniq.n` counts the number of unique elements in a vector or for each column in a matrix or data frame. By default, missing values are omitted before counting the number of unique elements.

Usage

```
uniq(data, ..., na.rm = TRUE, sort = TRUE, decreasing = FALSE, digits = NULL,
     table = TRUE, write = NULL, append = TRUE, check = TRUE, output = TRUE)

uniq.n(data, ..., na.rm = TRUE, digits = NULL, check = TRUE)
```

Arguments

<code>data</code>	a vector, factor, matrix, or data frame.
<code>...</code>	an expression indicating the variable names in data, e.g., <code>uniq(dat, x1, x2)</code> for selecting the variables <code>x1</code> and <code>x2</code> from the data frame <code>dat</code> . Note that the operators <code>+</code> , <code>-</code> , <code>~</code> , <code>:</code> , <code>::</code> , and <code>!</code> can also be used to select variables, see 'Details' in the df.subset function.
<code>na.rm</code>	logical: if TRUE (default), missing values are omitted before extracting unique elements. Note that missing values are always omitted when writing the output into an Excel file, i.e., <code>na.rm = TRUE</code> .
<code>sort</code>	logical: if TRUE (default), unique elements are sorted after extracting unique elements.
<code>decreasing</code>	logical: if TRUE, unique elements are sorted decreasing when specifying <code>sort = TRUE</code>
<code>digits</code>	an integer value indicating the number of decimal places to be used when rounding numeric values before extracting unique elements. By default, unique elements are extracted without rounding, i.e., <code>digits = NULL</code> .

table	logical: if TRUE (default), unique elements are printed in a data frame, if FALSE unique elements are printed in a list. Note that unique elements are always printed in a data frame when writing the output into an Excel file, i.e. table = TRUE.
write	a character string naming a file for writing the output into either a text file with file extension ".txt" (e.g., "Output.txt") or Excel file with file extension ".xlsx" (e.g., "Output.xlsx"). If the file name does not contain any file extension, an Excel file will be written.
append	logical: if TRUE (default), output will be appended to an existing text file with extension .txt specified in write, if FALSE existing text file will be overwritten.
check	logical: if TRUE (default), argument specification is checked.
output	logical: if TRUE (default), output is shown on the console.

Details

The function `uniqu` is a wrapper function in the form of `sort(unique(na.omit(x)))`, while the function `uniqu.n` is a wrapper function in the form of `length(unique(na.omit(x)))`.

Value

Returns an object of class `misty.object` when using the `uniqu` function, which is a list with following entries:

call	function call
type	type of analysis
data	a vector, factor, matrix, or data frame
args	specification of function arguments
result	list with unique elements

or a vector with the count number of unique elements for a vector, factor or each column in a matrix or data frame when using the `uniqu.n` function.

Author(s)

Takuya Yanagida

References

Becker, R. A., Chambers, J. M., & Wilks, A. R. (1988). *The New S Language*. Wadsworth & Brooks/Cole.

See Also

[df.duplicated](#), [df.unique](#)

Examples

```
#-----
# Extract Unique Elements, uniq() function

# Example 1a: Extract unique elements in a vector
uniq(airquality, Ozone)

# Example 1b: Extract unique elements in a vector, round elements
uniq(airquality, Wind, digits = 0)

# Example 1b: Extract unique elements in a vector, do not sort
uniq(airquality, Ozone, sort = FALSE)

# Example 1b: Extract unique elements in a vector, keep NA
uniq(airquality, Ozone, na.rm = FALSE)

# Example 2a: Extract unique elements in a data frame
uniq(airquality)

# Example 2a: Extract unique elements in list
uniq(airquality, table = FALSE)

#-----
# Count Number of Unique Elements, uniq.n() function

# Example 3a: Count number of unique elements in a vector
uniq.n(airquality, Ozone)

# Example 1b: Count number of unique elements for each variable in a data frame
uniq.n(airquality)
```

write.data

Write Data File in Table Format, SPSS, Excel, or Stata DTA File

Description

This function writes a (1) data file in CSV (.csv), DAT (.dat), or TXT (.txt) format using the `fwrite` function from the **data.table** package, (2) SPSS file (.sav) using the `write.sav` function, (3) Excel file (.xlsx) using the `write.xlsx` function, or a (4) Stata DTA file (.dta) using the `write.dta` function in the **misty** package. Note that the function `write.data` uses "," for decimal point and a semicolon ";" for the separator, while the function `write.data1` uses "." for decimal point and a comma "," for the separator when writing a CSV file.

Usage

```
write.data(x, file = "Data.csv", sep = ";", dec = ",", na = "",
           row.names = FALSE, col.names = TRUE, check = TRUE, ...)

write.data1(x, file = "Data.csv", sep = ",", dec = ".", na = "",
            row.names = FALSE, col.names = TRUE, check = TRUE, ...)
```

Arguments

x	a matrix or data frame to be written.
file	a character string indicating the name of the data file with the file extension .csv, .dat, .txt, .sav, .xlsx, or .dta. Note that the function will select an appropriate write-function depending on the file extension.
sep	a character string indicating the field separator, i.e., string for the delimiter. By default, the write.data function uses a semicolon ";", while the function write.data1 function uses a comma "," for writing a CSV, DAT, or TXT data file
dec	a character string indicating the decimal separator, i.e., string for decimal points. By default, the write.data function uses a comma ",", while the function write.data1 function uses a decimal point "." for writing a CSV, DAT, or TXT data file.
na	a character string to use for missing values in the data. By default, a blank string "" is used.
row.names	logical: if FALSE, row names are written.
col.names	logical: if TRUE (default), column names are written.
check	logical: if TRUE (default), argument specification is checked.
...	additional arguments to pass to the fwrite write.sav , write.xlsx , or write.dta function, see Arguments section in the help pages.

Details

Comma-Separated Values (CSV) File The function write.data writes CSV files based on the Excel convention for CSV files in some Western European locales by default, i.e., ";" as delimiter and "," for decimal points. Depending on the language setting of the operating system of the computer, the arguments sep and dec need to be specified to ";" and "." (see Example 1b). Alternatively, the function write.data1 that uses "," as delimiter and "." for decimal points by default can be used (see Example 1c).

Author(s)

Takuya Yanagida

References

- Barrett, T., Dowle, M., Srinivasan, A., Gorecki, J., Chirico, M., Hocking, T., & Schwendinger, B. (2024). data.table: Extension of 'data.frame'. R package version 1.16.0. <https://CRAN.R-project.org/package=data.table>
- Jeroen O. (2021). writexl: Export Data Frames to Excel 'xlsx' Format. R package version 1.4.0. <https://CRAN.R-project.org/package=writexl>
- Wickham H, Miller E, Smith D (2023). haven: Import and Export 'SPSS', 'Stata' and 'SAS' Files. R package version 2.5.3. <https://CRAN.R-project.org/package=haven>

See Also

[read.data](#), [read.sav](#), [write.sav](#), [write.xlsx](#), [read.dta](#), [write.dta](#), [read.mplus](#), [write.mplus](#)

Examples

```
## Not run:
# Example 1a: Write CSV data file, European format
write.data(mtcars, "European_CSV_Data.csv")

# Example 1b: Write CSV data file, American format
write.data(mtcars, "American_CSV_Data.csv", sep = ",", dec = ".")

# Example 1c: Write CSV data file, American format
write.data1(mtcars)

# Example 2: Write SPSS data file
write.data(mtcars, "SPSS_Data.sav")

# Example 3: Write Excel data file
write.data(mtcars, "Excel_Data.xlsx")

# Example 4: Write Stata data file
write.data(mtcars, "Stata_Data.dta")

## End(Not run)
```

write.dta

Write Stata DTA File

Description

This function writes a data frame or matrix into a Stata data file.

Usage

```
write.dta(x, file = "Stata_Data.dta", version = 14, label = NULL,
          str.thres = 2045, adjust.tz = TRUE, check = TRUE)
```

Arguments

x	a matrix or data frame to be written in Stata, vectors are coerced to a data frame.
file	a character string naming a file with or without file extension '.dta', e.g., "Stata_Data.dta" or "Stata_Data".
version	Stats file version to use. Supports versions 8-15.
label	dataset label to use, or NULL. Defaults to the value stored in the "label" attribute of data. Must be <= 80 characters.
str.thres	any character vector with a maximum length greater than str.thres bytes will be stored as a long string strL instead of a standard string str variable if version is greater or equal 13.
adjust.tz	this argument controls how the timezone of date-time values is treated when writing, see 'Details' in the write_dta function in the haven package.
check	logical: if TRUE (default), variable attributes specified in the argument var.attr is checked.

Note

This function is a modified copy of the `read_dta()` function in the **haven** package by Hadley Wickham, Evan Miller and Danny Smith (2023).

Author(s)

Hadley Wickham, Evan Miller and Danny Smith

References

Wickham H, Miller E, Smith D (2023). *haven: Import and Export 'SPSS', 'Stata' and 'SAS' Files*. R package version 2.5.3. <https://CRAN.R-project.org/package=haven>

See Also

[read.data](#), [write.data](#), [read.sav](#), [write.sav](#), [write.xlsx](#), [read.dta](#), [read.mplus](#), [write.mplus](#)

Examples

```
## Not run:

# Example 1: Write data frame 'mtcars' into the State data file 'mtcars.dta'
write.dta(mtcars, "mtcars.dta")

## End(Not run)
```

write.mplus

Write Mplus Data File

Description

This function writes a matrix or data frame to a tab-delimited file without variable names, a Mplus input template, and a text file with variable names. Note that only numeric variables are allowed, i.e., non-numeric variables will be removed from the data set. Missing data will be coded as a single numeric value.

Usage

```
write.mplus(x, file = "Mplus_Data.dat", data = TRUE, input = TRUE,
            var = FALSE, na = -99, check = TRUE)
```

Arguments

<code>x</code>	a matrix or data frame to be written to a tab-delimited file.
<code>file</code>	a character string naming a file with or without the file extension '.dat', e.g., "Mplus_Data.dat" or "Mplus_Data".
<code>data</code>	logical: if TRUE (default), Mplus data file is written in a text file named according to the argumentfile.

input	logical: if TRUE (default), Mplus input template is written in a text file named according to the argumentfile with the extension <code>_INPUT.inp</code> .
var	logical: if TRUE, variable names are written in a text file named according to the argumentfile with the extension <code>_VARNAMES.txt</code> .
na	a numeric value or character string representing missing values (NA) in the data set.
check	logical: if TRUE (default), argument specification is checked.

Value

Returns a character string indicating the variable names for the Mplus input file.

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

Muthen, L. K., & Muthen, B. O. (1998-2017). *Mplus User's Guide* (8th ed.). Muthen & Muthen.

See Also

[read.data](#), [write.data](#), [read.sav](#), [write.sav](#), [write.xlsx](#), [read.dta](#), [write.dta](#), [read.mplus](#)

Examples

```
## Not run:

# Example 1: Write Mplus Data File and a Mplus input template
write.mplus(mtcars)

# Example 2: Write Mplus Data File "mtcars.dat" and a Mplus input template "mtcars_INPUT.inp",
# missing values coded with -999,
# write variable names in a text file called "mtcars_VARNAMES.inp"
write.mplus(mtcars, file = "mtcars.dat", var = TRUE, na = -999)

## End(Not run)
```

write.result

Write Results of a misty Object into an Excel file

Description

This function writes the results of a `misty.object` into an Excel file.

Usage

```
write.result(x, file = "Results.xlsx", tri = x$args$tri,
            digits = x$args$digits, p.digits = x$args$p.digits,
            icc.digits = x$args$icc.digits, r.digits = x$args$r.digits,
            ess.digits = x$args$ess.digits, mcse.digits = x$args$mcse.digits,
            check = TRUE)
```

Arguments

<code>x</code>	misty object (<code>misty.object</code>) resulting from a misty function supported by the <code>write.result</code> function (see 'Details').
<code>file</code>	a character string naming a file with or without file extension '.xlsx', e.g., "Results.xlsx" or "Results".
<code>tri</code>	a character string or character vector indicating which triangular of the matrix to show on the console, i.e., both for upper and lower triangular, lower for the lower triangular, and upper for the upper triangular.
<code>digits</code>	an integer value indicating the number of decimal places digits to be used for displaying results.
<code>p.digits</code>	an integer indicating the number of decimal places to be used for displaying <i>p</i> -values.
<code>icc.digits</code>	an integer indicating the number of decimal places to be used for displaying intraclass correlation coefficients.
<code>r.digits</code>	an integer value indicating the number of decimal places to be used for displaying R-hat values.
<code>ess.digits</code>	an integer value indicating the number of decimal places to be used for displaying effective sample sizes.
<code>mcse.digits</code>	an integer value indicating the number of decimal places to be used for displaying Monte Carlo standard errors.
<code>check</code>	logical: if TRUE (default), argument specification is checked.

Details

Currently the function supports result objects from the following functions: `blimp.bayes`, `ci.cor`, `ci.mean`, `ci.median`, `ci.prop`, `ci.var`, `ci.sd`, `coeff.robust`, `coeff.std`, `cor.matrix`, `crosstab`, `descript`, `dominance.manual`, `dominance`, `effsize`, `freq`, `item.alpha`, `item.cfa`, `item.invar`, `item.omega`, `mplus.bayes`, `multilevel.cfa`, `multilevel.cor`, `multilevel.descript`, `multilevel.fit`, `multilevel.invar`, `multilevel.omega`, `na.auxiliary`, `na.coverage`, `na.descript`, `na.pattern`, `result.lca`, `summa` and `uniq`

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

Examples

```
## Not run:
#-----
# Example 1: item.cfa() function

# Load data set "HolzingerSwineford1939" in the lavaan package
data("HolzingerSwineford1939", package = "lavaan")

result <- item.cfa(HolzingerSwineford1939[, c("x1", "x2", "x3")], output = FALSE)
write.result(result, "CFA.xlsx")

#-----
# Example 2: multilevel.descript() function

# Load data set "Demo.twolevel" in the lavaan package
data("Demo.twolevel", package = "lavaan")

result <- multilevel.descript(y1:y3, data = Demo.twolevel, cluster = "cluster",
                             output = FALSE)
write.result(result, "Multilevel_Descript.xlsx")

## End(Not run)
```

write.sav

Write SPSS File

Description

This function writes a data frame or matrix into a SPSS file by either using the `write_sav()` function in the **haven** package by Hadley Wickham and Evan Miller (2019) or the free software *PSPP*.

Usage

```
write.sav(x, file = "SPSS_Data.sav", var.attr = NULL, pspp.path = NULL,
          digits = 2, write.csv = FALSE, sep = c(";", " "), na = "",
          write.sps = FALSE, check = TRUE)
```

Arguments

<code>x</code>	a matrix or data frame to be written in SPSS, vectors are coerced to a data frame.
<code>file</code>	a character string naming a file with or without file extension '.sav', e.g., "SPSS_Data.sav" or "SPSS_Data".
<code>var.attr</code>	a matrix or data frame with variable attributes used in the SPSS file, only 'variable labels' (column name label), 'value labels' column name values, and 'user-missing values' column name missing are supported (see 'Details').
<code>pspp.path</code>	a character string indicating the path where the PSPP folder is located on the computer, e.g.C:/Program Files/PSPP/.

digits	an integer value indicating the number of decimal places shown in the SPSS file for non-integer variables.
write.csv	logical: if TRUE, CSV file is written along with the SPSS file.
sep	a character string for specifying the CSV file, either ";" for the separator and "." for the decimal point (default, i.e. equivalent to write.csv2) or "," for the decimal point and ";" for the separator (i.e. equivalent to write.csv), must be one of both ";" (default) or ",".
na	a character string for specifying missing values in the CSV file.
write.sps	logical: if TRUE, SPSS syntax is written along with the SPSS file when using PSPP.
check	logical: if TRUE, variable attributes specified in the argument var.attr is checked.

Details

If arguments `pspp.path` is not specified (i.e., `pspp.path = NULL`), `write_sav()` function in the **haven** is used. Otherwise the object `x` is written as CSV file, which is subsequently imported into SPSS using the free software *PSPP* by executing a SPSS syntax written in R. Note that *PSPP* needs to be installed on your computer when using the `pspp.path` argument.

A SPSS file with 'variable labels', 'value labels', and 'user-missing values' is written by specifying the `var.attr` argument. Note that the number of rows in the matrix or data frame specified in `var.attr` needs to match with the number of columns in the data frame or matrix specified in `x`, i.e., each row in `var.attr` represents the variable attributes of the corresponding variable in `x`. In addition, column names of the matrix or data frame specified in `var.attr` needs to be labeled as label for 'variable labels, values for 'value labels', and missing for 'user-missing values'.

Labels for the values are defined in the column values of the matrix or data frame in `var.attr` using the equal-sign (e.g., `0 = female`) and are separated by a semicolon (e.g., `0 = female; 1 = male`).

User-missing values are defined in the column missing of the matrix or data frame in `var.attr`, either specifying one user-missing value (e.g., `-99`) or more than one but up to three user-missing values separated by a semicolon (e.g., `-77; -99`).

Note

Part of the function using *PSPP* was adapted from the `write.pspp()` function in the **miceadds** package by Alexander Robitzsch, Simon Grund and Thorsten Henke (2019).

Author(s)

Takuya Yanagida <takuya.yanagida@univie.ac.at>

References

- GNU Project (2018). *GNU PSPP for GNU/Linux* (Version 1.2.0). Boston, MA: Free Software Foundation. <https://www.gnu.org/software/pspp/>
- Wickham H., & Miller, E. (2019). *haven: Import and Export 'SPSS', 'Stata' and 'SAS' Files*. R package version 2.2.0.

Robitzsch, A., Grund, S., & Henke, T. (2019). *miceadds: Some additional multiple imputation functions, especially for mice*. R package version 3.4-17.

See Also

[read.data](#), [write.data](#), [read.sav](#), [write.xlsx](#), [read.dta](#), [write.dta](#), [read.mplus](#), [write.mplus](#)

Examples

```
## Not run:

dat <- data.frame(id = 1:5,
                  gender = c(NA, 0, 1, 1, 0),
                  age = c(16, 19, 17, NA, 16),
                  status = c(1, 2, 3, 1, 4),
                  score = c(511, 506, 497, 502, 491))

# Example 1: Write SPSS file using the haven package
write.sav(dat, file = "Dataframe_haven.sav")

# Example 2: Write SPSS file using PSPP,
# write CSV file and SPSS syntax along with the SPSS file
write.sav(dat, file = "Dataframe_PSPP.sav", pspp.path = "C:/Program Files/PSPP",
          write.csv = TRUE, write.sps = TRUE)

# Example 3: Specify variable attributes
# Note that it is recommended to manually specify the variables attributes in a CSV or
# Excel file which is subsequently read into R
attr <- data.frame(# Variable names
                   var = c("id", "gender", "age", "status", "score"),
                   # Variable labels
                   label = c("Identification number", "Gender", "Age in years",
                             "Migration background", "Achievement test score"),
                   # Value labels
                   values = c("", "0 = female; 1 = male", "",
                              "1 = Austria; 2 = former Yugoslavia; 3 = Turkey; 4 = other",
                              ""),
                   # User-missing values
                   missing = c("", "-99", "-99", "-99", "-99"))

# Example 4: Write SPSS file with variable attributes using the haven package
write.sav(dat, file = "Dataframe_haven_Attr.sav", var.attr = attr)

# Example 5: Write SPSS with variable attributes using PSPP
write.sav(dat, file = "Dataframe_PSPP_Attr.sav", var.attr = attr,
          pspp.path = "C:/Program Files/PSPP")

## End(Not run)
```

write.xlsxWrite Excel File

Description

This function calls the `write_excel()` function in the **writexl** package by Jeroen Ooms to write an Excel file (.xlsx).

Usage

```
write.xlsx(x, file = "Excel_Data.xlsx", col.names = TRUE, format = FALSE,  
           use.zip64 = FALSE, check = TRUE)
```

Arguments

<code>x</code>	a matrix, data frame or (named) list of matrices or data frames that will be written in the Excel file.
<code>file</code>	a character string naming a file with or without file extension '.xlsx', e.g., "My_Excel.xlsx" or "My_Excel".
<code>col.names</code>	logical: if TRUE, column names are written at the top of the Excel sheet.
<code>format</code>	logical: if TRUE, column names in the Excel file are centered and bold.
<code>use.zip64</code>	logical: if TRUE, zip64 to enable support for 4GB+ Excel files is used.
<code>check</code>	logical: if TRUE (default), argument specification is checked.

Details

This function supports strings, numbers, booleans, and dates.

Note

The function was adapted from the `write_excel()` function in the **writexl** package by Jeroen Ooms (2021).

Author(s)

Jeroen Ooms

References

Jeroen O. (2021). *writexl: Export Data Frames to Excel 'xlsx' Format*. R package version 1.4.0. <https://CRAN.R-project.org/package=writexl>

See Also

[read.data](#), [write.data](#), [read.sav](#), [write.sav](#), [read.dta](#), [write.dta](#), [read.mplus](#), [write.mplus](#)

Examples

```
## Not run:  
# Example 1: Write Excel file (.xlsx)  
write.xlsx(mtcars, file = "mtcars.xlsx")  
  
# Example 2: Write Excel file with multiple sheets (.xlsx)  
write.xlsx(list(cars = cars, mtcars = mtcars), file = "Excel_Sheets.xlsx")  
## End(Not run)
```

Index

- aov.b, [4](#), [12](#), [312–314](#), [353](#), [358](#), [362](#), [366](#)
- aov.w, [7](#), [8](#), [88](#), [312–315](#), [358](#), [366](#)
- as.na, [294](#), [296](#), [298](#), [300](#), [302](#), [304](#), [311](#)
- as.na (na.as), [289](#)

- blimp, [13](#), [21](#), [26](#), [30](#), [33](#)
- blimp.bayes, [15](#), [17](#), [26](#), [30](#), [33](#), [36](#), [375](#)
- blimp.plot, [15](#), [21](#), [22](#), [26](#), [30](#), [33](#), [36](#)
- blimp.print, [14](#), [15](#), [21](#), [26](#), [28](#), [33](#), [35](#), [36](#)
- blimp.run, [15](#), [21](#), [26](#), [30](#), [31](#), [36](#)
- blimp.update, [15](#), [21](#), [26](#), [30](#), [33](#), [33](#)

- center, [37](#), [204](#)
- cfa.satcor (na.satcor), [304](#)
- check.collin, [42](#), [46](#), [50](#), [347](#), [348](#)
- check.outlier, [44](#), [45](#), [50](#)
- check.resid, [47](#)
- chr.color, [51](#), [54–56](#), [58](#), [59](#), [140](#), [144](#)
- chr.grep, [52](#), [53](#), [55](#), [56](#), [58](#), [59](#)
- chr.grepl, [52](#), [54–56](#), [58](#), [59](#)
- chr.grepl (chr.grep), [53](#)
- chr.gsub, [52](#), [54](#), [54](#), [56](#), [58](#), [59](#)
- chr.omit, [52](#), [54](#), [55](#), [56](#), [58](#), [59](#)
- chr.trim, [52](#), [54–56](#), [57](#), [59](#)
- chr.trunc, [52](#), [54–56](#), [58](#), [58](#)
- ci.cor, [59](#), [75](#), [80](#), [89](#), [94](#), [100](#), [105](#), [312–314](#), [375](#)
- ci.mean, [7](#), [12](#), [70](#), [73](#), [84](#), [98](#), [105](#), [138](#), [358](#), [362](#), [366](#), [375](#)
- ci.mean.diff, [7](#), [12](#), [70](#), [80](#), [82](#), [88](#), [98](#), [105](#), [138](#), [358](#), [362](#), [366](#)
- ci.mean.w, [85](#)
- ci.median, [84](#), [88](#), [94](#), [98](#), [105](#), [138](#), [375](#)
- ci.median (ci.mean), [73](#)
- ci.prop, [70](#), [80](#), [84](#), [88](#), [88](#), [94](#), [98](#), [105](#), [138](#), [375](#)
- ci.prop.diff, [94](#), [96](#), [105](#), [138](#)
- ci.sd, [70](#), [80](#), [84](#), [88](#), [94](#), [98](#), [138](#), [375](#)
- ci.sd (ci.var), [99](#)
- ci.var, [70](#), [80](#), [84](#), [88](#), [94](#), [98](#), [99](#), [138](#), [375](#)

- clear, [107](#)
- cluster.rwg, [41](#), [108](#)
- cluster.scores, [41](#), [110](#), [110](#), [201](#), [251](#)
- coding, [41](#), [112](#), [204](#), [326](#)
- coeff.robust, [115](#), [334](#), [347](#), [348](#), [375](#)
- coeff.std, [119](#), [120](#), [161](#), [162](#), [346](#), [348](#), [375](#)
- cohens.d, [7](#), [12](#), [125](#), [133](#), [167](#), [355](#), [358](#), [360](#), [362](#), [366](#)
- cor.matrix, [70](#), [128](#), [130](#), [167](#), [342](#), [346](#), [348](#), [375](#)
- cor.test, [132](#), [342](#)
- crosstab, [134](#), [138](#), [169](#), [375](#)

- descript, [80](#), [84](#), [88](#), [94](#), [98](#), [105](#), [135](#), [136](#), [169](#), [256](#), [345](#), [348](#), [375](#)
- df.check, [139](#), [142](#), [144](#), [147](#), [149](#), [151–153](#), [155](#), [157](#)
- df.duplicated, [140](#), [141](#), [144](#), [147](#), [149](#), [151–153](#), [155](#), [157](#), [369](#)
- df.head, [140](#), [142](#), [143](#), [147](#), [149](#), [151–153](#), [155](#), [157](#)
- df.long, [140](#), [142](#), [144](#), [145](#), [149](#), [151–153](#), [155](#), [157](#)
- df.merge, [140](#), [142](#), [144](#), [147](#), [148](#), [151–153](#), [155](#), [157](#)
- df.move, [140](#), [142](#), [144](#), [147](#), [149](#), [150](#), [152](#), [153](#), [155](#), [157](#)
- df.rbind, [140](#), [142](#), [144](#), [147](#), [149](#), [151](#), [151](#), [153](#), [155](#), [157](#)
- df.rename, [140](#), [142](#), [144](#), [147](#), [149](#), [151](#), [152](#), [153](#), [155](#), [157](#)
- df.sort, [140](#), [142](#), [144](#), [147](#), [149](#), [151–153](#), [154](#), [157](#)
- df.subset, [38](#), [60](#), [74](#), [86](#), [89](#), [100](#), [108](#), [111](#), [131](#), [134](#), [137](#), [140–142](#), [144](#), [145](#), [147](#), [149–153](#), [155](#), [155](#), [165](#), [168](#), [175](#), [178](#), [188](#), [195](#), [198](#), [200](#), [202](#), [242](#), [248](#), [253](#), [260](#), [268](#), [273](#), [289](#), [292](#), [295](#), [297](#), [299](#), [301](#), [303](#), [306](#), [325](#), [344](#), [368](#)

- df.tail, [140](#), [142](#), [147](#), [149](#), [151–153](#), [155](#), [157](#)
- df.tail (df.head), [143](#)
- df.unique, [140](#), [144](#), [147](#), [149](#), [151–153](#), [155](#), [157](#), [369](#)
- df.unique (df.duplicated), [141](#)
- df.wide, [140](#), [142](#), [144](#), [149](#), [151–153](#), [155](#), [157](#)
- df.wide (df.long), [145](#)
- dominance, [158](#), [162](#), [375](#)
- dominance.manual, [161](#), [161](#), [375](#)
- effsize, [128](#), [133](#), [165](#), [375](#)
- format.POSIXct, [336](#)
- freq, [135](#), [138](#), [168](#), [315](#), [375](#)
- growth.satcor (na.satcor), [304](#)
- indirect, [170](#), [265–267](#)
- item.alpha, [174](#), [185](#), [199](#), [201](#), [375](#)
- item.cfa, [175](#), [177](#), [178](#), [192](#), [195](#), [197](#), [201](#), [245](#), [292](#), [375](#)
- item.invar, [177](#), [187](#), [197](#), [375](#)
- item.omega, [177](#), [185](#), [194](#), [197](#), [199](#), [201](#), [275](#), [375](#)
- item.reverse, [41](#), [114](#), [177](#), [197](#), [198](#), [204](#), [326](#)
- item.scores, [41](#), [111](#), [177](#), [185](#), [197](#), [199](#), [200](#)
- kurtosis (skewness), [343](#)
- lagged, [202](#)
- lavaan.satcor (na.satcor), [304](#)
- libraries, [205](#)
- library, [205](#)
- lm, [44](#), [46](#)
- mplus, [206](#), [214](#), [218](#), [223](#), [234](#), [236](#), [240](#)
- mplus.bayes, [208](#), [209](#), [218](#), [223](#), [234](#), [236](#), [240](#), [375](#)
- mplus.lca, [208](#), [214](#), [214](#), [223](#), [234](#), [236](#), [240](#), [331](#)
- mplus.plot, [208](#), [214](#), [218](#), [219](#), [234](#), [236](#), [240](#)
- mplus.print, [207](#), [208](#), [214](#), [218](#), [223](#), [229](#), [236](#), [238](#), [240](#)
- mplus.run, [208](#), [214](#), [218](#), [223](#), [234](#), [235](#), [240](#), [331](#)
- mplus.update, [208](#), [214](#), [218](#), [223](#), [234](#), [236](#), [237](#)
- multilevel.cfa, [241](#), [259](#), [263](#), [271](#), [275](#), [375](#)
- multilevel.cor, [245](#), [248](#), [259](#), [263](#), [271](#), [275](#), [283](#), [288](#), [346](#), [375](#)
- multilevel.descript, [111](#), [135](#), [138](#), [169](#), [245](#), [251](#), [253](#), [259](#), [263](#), [271](#), [275](#), [283](#), [288](#), [375](#)
- multilevel.fit, [245](#), [257](#), [271](#), [275](#), [375](#)
- multilevel.icc, [111](#), [133](#), [251](#), [254–256](#), [260](#), [283](#), [288](#)
- multilevel.indirect, [174](#), [264](#), [283](#), [288](#)
- multilevel.invar, [192](#), [245](#), [259](#), [267](#), [275](#), [375](#)
- multilevel.omega, [245](#), [259](#), [271](#), [273](#), [375](#)
- multilevel.r2, [276](#), [287](#), [288](#), [313](#), [346](#)
- multilevel.r2.manual, [285](#)
- na.as, [289](#), [294](#), [296](#), [298](#), [300](#), [302](#), [304](#), [311](#)
- na.auxiliary, [128](#), [133](#), [290](#), [291](#), [296](#), [298](#), [300](#), [302](#), [304](#), [311](#), [375](#)
- na.coverage, [290](#), [294](#), [295](#), [298](#), [300](#), [302](#), [304](#), [311](#), [375](#)
- na.descript, [135](#), [138](#), [169](#), [290](#), [294](#), [296](#), [297](#), [300](#), [302](#), [304](#), [311](#), [375](#)
- na.indicator, [290](#), [294](#), [296](#), [298](#), [299](#), [302](#), [304](#), [311](#)
- na.pattern, [290](#), [294](#), [296](#), [298](#), [300](#), [300](#), [304](#), [311–314](#), [375](#)
- na.prop, [290](#), [294](#), [296](#), [298](#), [300](#), [302](#), [303](#), [311](#)
- na.satcor, [304](#)
- na.test, [290](#), [294](#), [296](#), [298](#), [300](#), [302](#), [304](#), [306](#)
- p.adjust, [9](#), [131](#), [249](#)
- plot.misty.object, [311](#)
- print.misty.object, [315](#)
- prop.test, [342](#)
- rbind, [152](#)
- read.data, [316](#), [319](#), [321](#), [322](#), [324](#), [371](#), [373](#), [374](#), [378](#), [379](#)
- read.dta, [318](#), [318](#), [321](#), [322](#), [324](#), [371](#), [373](#), [374](#), [378](#), [379](#)
- read.mplus, [208](#), [214](#), [218](#), [223](#), [234](#), [236](#), [240](#), [318](#), [319](#), [320](#), [322](#), [324](#), [331](#), [371](#), [373](#), [374](#), [378](#), [379](#)
- read.sav, [318](#), [319](#), [321](#), [321](#), [324](#), [371](#), [373](#), [374](#), [378](#), [379](#)
- read.xlsx, [323](#)

rec, [41](#), [114](#), [199](#), [204](#), [325](#)
require, [205](#)
restart, [108](#), [328](#)
result.lca, [328](#), [375](#)
robust.lmer, [333](#)

script.close, [336](#), [337](#), [340](#)
script.close (script.open), [338](#)
script.copy, [335](#), [337](#), [339](#)
script.new, [336](#), [337](#), [340](#)
script.open, [336](#), [337](#), [338](#), [340](#)
script.save, [336](#), [337](#), [339](#), [340](#)
script.save (script.open), [338](#)
sem.satcor (na.satcor), [304](#)
setsource, [108](#), [336](#), [337](#), [339](#), [340](#)
size.cor, [133](#)
size.cor (size.mean), [341](#)
size.mean, [341](#)
size.prop (size.mean), [341](#)
skewness, [343](#)
summa, [334](#), [346](#), [375](#)

test.levene, [7](#), [312–314](#), [351](#), [358](#), [362](#)
test.t, [7](#), [12](#), [80](#), [84](#), [88](#), [128](#), [312–315](#), [342](#),
 [353](#), [354](#), [362](#), [366](#)
test.welch, [7](#), [312–315](#), [353](#), [358](#), [359](#), [366](#)
test.z, [7](#), [12](#), [80](#), [84](#), [88](#), [128](#), [315](#), [358](#), [362](#),
 [363](#)

uniq, [368](#), [375](#)

write.data, [318](#), [319](#), [321](#), [322](#), [324](#), [370](#),
 [373](#), [374](#), [378](#), [379](#)
write.data1 (write.data), [370](#)
write.dta, [318](#), [319](#), [321](#), [322](#), [324](#), [371](#), [372](#),
 [374](#), [378](#), [379](#)
write.mplus, [208](#), [214](#), [218](#), [223](#), [234](#), [236](#),
 [240](#), [318](#), [319](#), [321](#), [322](#), [324](#), [331](#),
 [371](#), [373](#), [373](#), [378](#), [379](#)
write.result, [119](#), [133](#), [135](#), [161](#), [162](#), [169](#),
 [177](#), [192](#), [197](#), [251](#), [256](#), [296](#), [298](#),
 [302](#), [374](#)
write.sav, [318](#), [319](#), [321](#), [322](#), [324](#), [371](#), [373](#),
 [374](#), [376](#), [379](#)
write.xlsx, [318](#), [319](#), [321](#), [322](#), [324](#), [371](#),
 [373](#), [374](#), [378](#), [379](#)