

Package ‘mirt’

July 22, 2025

Version 1.44.0

Type Package

Title Multidimensional Item Response Theory

Description Analysis of discrete response data using
unidimensional and multidimensional item analysis models under the Item
Response Theory paradigm (Chalmers (2012) <[doi:10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)>).
Exploratory and confirmatory item factor analysis models
are estimated with quadrature (EM) or stochastic (MHRM) methods. Confirmatory
bi-factor and two-tier models are available for modeling item testlets using
dimension reduction EM algorithms, while multiple group analyses and
mixed effects designs are included for detecting differential item, bundle,
and test functioning, and for modeling item and person covariates.
Finally, latent class models such as the DINA, DINO, multidimensional latent class,
mixture IRT models, and zero-inflated response models are supported, as well
as a wide family of probabilistic unfolding models.

VignetteBuilder knitr

Depends stats, R (>= 4.1.0), stats4, lattice, methods

Imports GPArotation, gridExtra, Matrix (>= 1.5-0), Rcpp, mgcv, vegan,
Deriv, splines, pbapply (>= 1.3-0), dcurver, SimDesign

Suggests boot, latticeExtra, directlabels, shiny, knitr, markdown,
Rsolnp, nloptr, sirt, plink, mirtCAT, testthat (>= 3.0.0)

ByteCompile yes

LazyLoad yes

LazyData yes

LinkingTo Rcpp, RcppArmadillo

License GPL (>= 3)

Encoding UTF-8

Repository CRAN

Maintainer Phil Chalmers <rphilip.chalmers@gmail.com>

URL <https://philchalmers.github.io/mirt/>,
<https://github.com/philchalmers/mirt/wiki>,
<https://groups.google.com/forum/#!forum/mirt-package>

BugReports <https://github.com/philchalmers/mirt/issues?state=open>

RoxygenNote 7.3.2

Config/testthat/edition 3

NeedsCompilation yes

Author Phil Chalmers [aut, cre] (ORCID:
[<https://orcid.org/0000-0001-5332-2810>](https://orcid.org/0000-0001-5332-2810)),
Joshua Pritikin [ctb],
Alexander Robitzsch [ctb],
Mateusz Zoltak [ctb],
KwonHyun Kim [ctb],
Carl F. Falk [ctb],
Adam Meade [ctb],
Lennart Schneider [ctb],
David King [ctb],
Chen-Wei Liu [ctb],
Ogreden Oguzhan [ctb]

Date/Publication 2025-01-28 21:50:02 UTC

Contents

mirt-package	4
anova-method	5
areainfo	7
ASVAB	8
Attitude	9
averageMI	10
bfactor	12
Bock1997	18
boot.LR	18
boot.mirt	20
coef-method	21
createGroup	23
createItem	25
deAyala	29
DeltaMethod	29
DIF	30
DiscreteClass-class	36
draw_parameters	37
DRF	38
DTF	45
empirical_ES	48
empirical_plot	51

empirical_rxx	53
estfun.AllModelClass	54
expand.table	56
expected.item	58
expected.test	59
extract.group	61
extract.item	62
extract.mirt	63
fixedCalib	65
fixef	68
fscores	70
gen.difficulty	76
imputeMissing	77
itemfit	79
itemGAM	84
iteminfo	88
itemplot	89
itemstats	92
key2binary	94
lagrange	95
likert2int	97
logLik-method	99
LSAT6	100
LSAT7	101
M2	102
marginal_rxx	104
MDIFF	105
mdirt	106
MDISC	113
mirt	114
mirt.model	143
mirtCluster	148
MixedClass-class	149
mixedmirt	150
MixtureClass-class	158
mod2values	159
multipleGroup	160
MultipleGroupClass-class	171
numerical_deriv	172
personfit	173
PLCI.mirt	175
plot,MultipleGroupClass,missing-method	177
poly2dich	181
print-method	182
print.mirt_df	182
print.mirt_list	183
print.mirt_matrix	183
probtrace	184

randef	185
RCI	186
read.mirt	189
remap.distance	191
residuals-method	192
reverse.score	195
RMSD_DIF	197
SAT12	200
Science	201
secondOrderTest	202
show-method	203
SIBTEST	204
simdata	209
SingleGroupClass-class	217
SLF	218
summary-method	219
testinfo	220
thetaComb	222
traditional2mirt	223
vcov-method	224
wald	225
Index	228

mirt-package

Full information maximum likelihood estimation of IRT models.

Description

Full information maximum likelihood estimation of multidimensional IRT models

Details

Analysis of dichotomous and polytomous response data using unidimensional and multidimensional latent trait models under the Item Response Theory (IRT) paradigm. Exploratory and confirmatory models can be estimated with quadrature (EM) or stochastic (MHRM) methods. Confirmatory bi-factor and two-tier analyses are available for modeling item testlets. Multiple group analysis and mixed effects designs also are available for detecting differential item and test functioning as well as modeling item and person covariates. Finally, latent class models such as the DINA, DINO, multidimensional latent class, mixture and zero-inflated IRT models, and several other discrete variable models are supported.

Users interested in the most recent version of this package can visit <https://github.com/philchalmers/mirt> and follow the instructions for installing the package from source. Questions regarding the package can be sent to the mirt-package Google Group, located at <https://groups.google.com/forum/#!forum/mirt-package>. User contributed files, workshop files, and evaluated help files are also available on the package wiki (<https://github.com/philchalmers/mirt/wiki>).

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

 anova-method

Compare nested models with likelihood-based statistics

Description

Compare nested models using likelihood ratio test (X2), Akaike Information Criterion (AIC), Bayesian Information Criterion (BIC), Sample-Size Adjusted BIC (SABIC), and Hannan-Quinn (HQ) Criterion. When given a sequence of objects, anova tests the models against one another in the order specified. Note that the object inputs should be ordered in terms of most constrained model to least constrained.

Usage

```
## S4 method for signature 'SingleGroupClass'
anova(
  object,
  object2,
  ...,
  bounded = FALSE,
  mix = 0.5,
  frame = 1,
  verbose = FALSE
)
```

Arguments

object	an object of class SingleGroupClass, MultipleGroupClass, or MixedClass, reflecting the most constrained model fitted
object2	a second model estimated from any of the mirt package estimation methods
...	additional less constrained model objects to be compared sequentially to the previous model
bounded	logical; are the two models comparing a bounded parameter (e.g., comparing a single 2PL and 3PL model with 1 df)? If TRUE then a 50:50 mix of chi-squared distributions is used to obtain the p-value
mix	proportion of chi-squared mixtures. Default is 0.5
frame	(internal parameter not for standard use)
verbose	(deprecated argument)

Value

a data.frame/mirt_df object

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```
## Not run:
x <- mirt(Science, 1)
x2 <- mirt(Science, 2)
anova(x, x2)

# compare three models sequentially (X2 not always meaningful)
x3 <- mirt(Science, 1, 'gpcm')
x4 <- mirt(Science, 1, 'nominal')
anova(x, x2, x3, x4)

# in isolation
anova(x)

# with priors on first model
model <- "Theta = 1-4
        PRIOR = (1-4, a1, lnorm, 0, 10)"
xp <- mirt(Science, model)
anova(xp, x2)
anova(xp)

# bounded parameter
dat <- expand.table(LSAT7)
mod <- mirt(dat, 1)
mod2 <- mirt(dat, 1, itemtype = c(rep('2PL', 4), '3PL'))
anova(mod, mod2) #unbounded test
anova(mod, mod2, bounded = TRUE) #bounded

# priors
model <- 'F = 1-5
        PRIOR = (5, g, norm, -1, 1)'
mod1b <- mirt(dat, model, itemtype = c(rep('2PL', 4), '3PL'))
anova(mod1b)

model2 <- 'F = 1-5
        PRIOR = (1-5, g, norm, -1, 1)'
mod2b <- mirt(dat, model2, itemtype = '3PL')
anova(mod1b, mod2b)

## End(Not run)
```

areainfo	<i>Function to calculate the area under a selection of information curves</i>
----------	---

Description

Compute the area of a test or item information function over a definite integral range.

Usage

```
areainfo(  
  x,  
  theta_lim,  
  which.items = 1:extract.mirt(x, "nitems"),  
  group = NULL,  
  ...  
)
```

Arguments

x	an object of class 'SingleGroupClass', or an object of class 'MultipleGroup-Class' if a suitable group input were supplied
theta_lim	range of integration to be computed
which.items	an integer vector indicating which items to include in the expected information function. Default uses all possible items
group	group argument to pass to extract.group function. Required when the input object is a multiple-group model
...	additional arguments passed to integrate

Value

a data.frame with the lower and upper integration range, the information area within the range (Info), the information area over the range -10 to 10 (Total.Info), proportion of total information given the integration range (Info.Proportion), and the number of items included (nitems)

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```

dat <- expand.table(LSAT7)
mod <- mirt(dat, 1)

areainfo(mod, c(-2,0), which.items = 1) #item 1
## Not run:
areainfo(mod, c(-2,0), which.items = 1:3) #items 1 to 3
areainfo(mod, c(-2,0)) # all items (total test information)

# plot the area
area <- areainfo(mod, c(-2,0))
Theta <- matrix(seq(-3,3, length.out=1000))
info <- testinfo(mod, Theta)
plot(info ~ Theta, type = 'l')

pick <- Theta >= -2 & Theta <= 0
polygon(c(-2, Theta[pick], 0), c(0, info[pick], 0), col='lightblue')
text(x = 2, y = 0.5, labels = paste("Total Information:", round(area$TotalInfo, 3),
  "\n\nInformation in (-2, 0):", round(area$Info, 3),
  paste("(", round(100 * area$Proportion, 2), "%)", sep = " ")), cex = 1.2)

## End(Not run)

```

ASVAB

Description of ASVAB data

Description

Table of counts extracted from Mislvey (1985). Data the 16 possible response patterns observed for four items from the arithmetic reasoning test of the Armed Services Vocational Aptitude Battery (ASVAB), Form 8A, from samples of white males and females and black males and females.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Mislevy, R. J. (1985). Estimation of latent group effects. *Journal of the American Statistical Association*, 80, 993-997.

Examples

```

data(ASVAB)
datWM <- expand.table(subset(ASVAB, select=c(Item.1:Item.4, White_Male)))
datWF <- expand.table(subset(ASVAB, select=c(Item.1:Item.4, White_Female)))
datBM <- expand.table(subset(ASVAB, select=c(Item.1:Item.4, Black_Male)))
datBF <- expand.table(subset(ASVAB, select=c(Item.1:Item.4, Black_Female)))

```

```

dat <- rbind(datWM, datWF, datBM, datBF)
sex <- rep(c("Male", "Female", "Male", "Female"),
  times=c(nrow(datWM), nrow(datWF), nrow(datBM), nrow(datBF))) |> factor()
color <- rep(c("White", "Black"),
  times=c(nrow(datWM) + nrow(datWF), nrow(datBM) + nrow(datBF))) |> factor()
group <- sex:color

itemstats(dat, group=group)

```

Attitude	<i>Description of Attitude data</i>
----------	-------------------------------------

Description

Table of counts extracted from Andrich (1988). Data the response patterns observed for an eight item survey.

Details

The items in this survey were:

1. Capital punishment is one of the most hideous practices of our time.
2. The state cannot teach the sacredness of human life by destroying it.
3. Capital punishment is not an effective deterrent to crime.
4. I don't believe in capital punishment but I am not sure it isn't necessary.
5. I think capital punishment is necessary but I wish it were not.
6. Until we find a more civilized way to prevent crime we must have capital punishment.
7. Capital punishment is justified because it does act as a deterrent to crime.
8. Capital punishment gives the criminal what he deserves.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Andrich, D. (1988). The Application of an Unfolding Model of the PIRT Type to the Measurement of Attitude. *Applied Psychological Measurement*, 12, 33-51.

Examples

```

head(Attitude)
df <- expand.table(Attitude)
itemstats(df)

## Not run:

# estimate SSLM with estimated " latitude of acceptance" (rho)
mod.rho <- mirt(df, 1, itemtype = 'sslm')
coef(mod.rho)
coef(mod.rho, simplify=TRUE) # slope-intercept-log_rho
coef(mod.rho, simplify=TRUE, IRTpars=TRUE) # discrimination-difficulty-rho
plot(mod.rho)
plot(mod.rho, type = 'trace')

# without estimating rho, and fixing to rho^2 = 1 (hence,
# log_rho = -exp(1) = -2.718282 in order to obtain (exp(exp(log_rho))) = 1)
syntax <- "Theta = 1-8
          FIXED = (1-8, log_rho1)
          START = (1-8, log_rho1, -2.71828)"
mod <- mirt(df, syntax, itemtype = 'sslm') # model found in Andrich (1988)
coef(mod)
coef(mod, simplify=TRUE) # slope-intercept-log_rho
coef(mod, simplify=TRUE, IRTpars=TRUE) # discrimination-difficulty-rho
plot(mod)
plot(mod, type = 'trace') # notice that all curves have a fixed height of .5

# goodness of fit (less constrained model fits better)
anova(mod, mod.rho) # original model fits much worse
M2(mod)
M2(mod.rho)
itemfit(mod, p.adjust='fdr')
itemfit(mod.rho, p.adjust='fdr')

## End(Not run)

```

averageMI

Collapse values from multiple imputation draws

Description

This function computes updated parameter and standard error estimates using multiple imputation methodology. Given a set of parameter estimates and their associated standard errors the function returns the weighted average of the overall between and within variability due to the multiple imputations according to Rubin's (1987) methodology.

Usage

```
averageMI(par, SEpar, as.data.frame = TRUE)
```

Arguments

`par` a list containing parameter estimates which were computed the imputed datasets
`SEpar` a list containing standard errors associated with `par`
`as.data.frame` logical; return a data.frame instead of a list? Default is TRUE

Value

returns a list or data.frame containing the updated averaged parameter estimates, standard errors, and t-values with the associated degrees of freedom and two tailed p-values

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
 Rubin, D.B. (1987) Multiple Imputation for Nonresponse in Surveys. Wiley & Sons, New York.

Examples

```
## Not run:

# simulate data
set.seed(1234)
N <- 1000

# covariates
X1 <- rnorm(N); X2 <- rnorm(N)
covdata <- data.frame(X1, X2)
Theta <- matrix(0.5 * X1 + -1 * X2 + rnorm(N, sd = 0.5))

# items and response data
a <- matrix(1, 20); d <- matrix(rnorm(20))
dat <- simdata(a, d, 1000, itemtype = '2PL', Theta=Theta)

mod1 <- mirt(dat, 1, 'Rasch', covdata=covdata, formula = ~ X1 + X2)
coef(mod1, simplify=TRUE)

# draw plausible values for secondary analyses
pv <- fscores(mod1, plausible.draws = 10)
pvmods <- lapply(pv, function(x, covdata) lm(x ~ covdata$X1 + covdata$X2),
  covdata=covdata)

# compute Rubin's multiple imputation average
```

```

so <- lapply(pvmods, summary)
par <- lapply(so, function(x) x$coefficients[, 'Estimate'])
SEpar <- lapply(so, function(x) x$coefficients[, 'Std. Error'])
averageMI(par, SEpar)

## End(Not run)

```

bfactor

Full-Information Item Bi-factor and Two-Tier Analysis

Description

bfactor fits a confirmatory maximum likelihood two-tier/bifactor/testlet model to dichotomous and polytomous data under the item response theory paradigm. The IRT models are fit using a dimensional reduction EM algorithm so that regardless of the number of specific factors estimated the model only uses the number of factors in the second-tier structure plus 1. For the bifactor model the maximum number of dimensions is only 2 since the second-tier only consists of a ubiquitous unidimensional factor. See [mirt](#) for appropriate methods to be used on the objects returned from the estimation.

Usage

```

bfactor(
  data,
  model,
  model2 = paste0("G = 1-", ncol(data)),
  group = NULL,
  quadpts = NULL,
  invariance = "",
  ...
)

```

Arguments

data	a matrix or data.frame that consists of numerically ordered data, organized in the form of integers, with missing data coded as NA
model	a numeric vector specifying which factor loads on which item. For example, if for a 4 item test with two specific factors, the first specific factor loads on the first two items and the second specific factor on the last two, then the vector is <code>c(1, 1, 2, 2)</code> . For items that should only load on the second-tier factors (have no specific component) NA values may be used as place-holders. These numbers will be translated into a format suitable for <code>mirt.model()</code> , combined with the definition in <code>model2</code> , with the letter 'S' added to the respective factor number. Alternatively, input can be specified using the <code>mirt.model</code> syntax with the restriction that each item must load on exactly one specific factor (or no specific factors, if it is only predicted by the general factor specified in <code>model2</code>)

model2	a two-tier model specification object defined by <code>mirt.model()</code> or a string to be passed to <code>mirt.model</code> . By default the model will fit a unidimensional model in the second-tier, and therefore be equivalent to the bifactor model
group	a factor variable indicating group membership used for multiple group analyses
quadpts	number of quadrature nodes to use after accounting for the reduced number of dimensions. Scheme is the same as the one used in <code>mirt</code> , however it is in regards to the reduced dimensions (e.g., a bifactor model has 2 dimensions to be integrated)
invariance	see <code>multipleGroup</code> for details, however, the specific factor variances and means will be constrained according to the dimensional reduction algorithm
...	additional arguments to be passed to the estimation engine. See <code>mirt</code> for more details and examples

Details

`bfactor` follows the item factor analysis strategy explicated by Gibbons and Hedeker (1992), Gibbons et al. (2007), and Cai (2010). Nested models may be compared via an approximate chi-squared difference test or by a reduction in AIC or BIC (accessible via `anova`). See `mirt` for more details regarding the IRT estimation approach used in this package.

The two-tier model has a specific block diagonal covariance structure between the primary and secondary latent traits. Namely, the secondary latent traits are assumed to be orthogonal to all traits and have a fixed variance of 1, while the primary traits can be organized to vary and covary with other primary traits in the model.

$$\Sigma_{two-tier} = \begin{pmatrix} G & 0 \\ 0 & diag(S) \end{pmatrix}$$

The bifactor model is a special case of the two-tier model when G above is a 1x1 matrix, and therefore only 1 primary factor is being modeled. Evaluation of the numerical integrals for the two-tier model requires only $ncol(G) + 1$ dimensions for integration since the S second order (or 'specific') factors require only 1 integration grid due to the dimension reduction technique.

Note: for multiple group two-tier analyses only the second-tier means and variances should be freed since the specific factors are not treated independently due to the dimension reduction technique.

Value

function returns an object of class `SingleGroupClass` (`SingleGroupClass-class`) or `MultipleGroupClass` (`MultipleGroupClass-class`).

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Cai, L. (2010). A two-tier full-information item factor analysis model with applications. *Psychometrika*, 75, 581-612.

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Bradlow, E.T., Wainer, H., & Wang, X. (1999). A Bayesian random effects model for testlets. *Psychometrika*, 64, 153-168.

Gibbons, R. D., & Hedeker, D. R. (1992). Full-information Item Bi-Factor Analysis. *Psychometrika*, 57, 423-436.

Gibbons, R. D., Darrell, R. B., Hedeker, D., Weiss, D. J., Segawa, E., Bhaumik, D. K., Kupfer, D. J., Frank, E., Grochocinski, V. J., & Stover, A. (2007). Full-Information item bifactor analysis of graded response data. *Applied Psychological Measurement*, 31, 4-19.

Wainer, H., Bradlow, E.T., & Wang, X. (2007). Testlet response theory and its applications. New York, NY: Cambridge University Press.

See Also

[mirt](#)

Examples

```
## Not run:

### load SAT12 and compute bifactor model with 3 specific factors
data(SAT12)
data <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
specific <- c(2,3,2,3,3,2,1,2,1,1,1,3,1,3,1,2,1,1,3,3,1,1,3,1,3,3,1,3,2,3,1,2)
mod1 <- bfactor(data, specific)
summary(mod1)
itemplot(mod1, 18, drop.zeros = TRUE) #drop the zero slopes to allow plotting

# alternative model definition via ?mirt.model syntax
specific2 <- "S1 = 7,9,10,11,13,15,17,18,21,22,24,27,31
             S2 = 1,3,6,8,16,29,32
             S3 = 2,4,5,12,14,19,20,23,25,26,28,30"
mod2 <- bfactor(data, specific2)
anova(mod1, mod2) # same

# also equivalent using item names instead (not run)
specific3 <- "S1 = Item.7, Item.9, Item.10, Item.11, Item.13, Item.15,
             Item.17, Item.18, Item.21, Item.22, Item.24, Item.27, Item.31
             S2 = Item.1, Item.3, Item.6, Item.8, Item.16, Item.29, Item.32
             S3 = Item.2, Item.4, Item.5, Item.12, Item.14, Item.19,
             Item.20, Item.23, Item.25, Item.26, Item.28, Item.30"
# mod3 <- bfactor(data, specific3)
# anova(mod1, mod2, mod3) # all same

### Try with fixed guessing parameters added
guess <- rep(.1,32)
mod2 <- bfactor(data, specific, guess = guess)
coef(mod2)
anova(mod1, mod2)
```

```
## don't estimate specific factor for item 32
specific[32] <- NA
mod3 <- bfactor(data, specific)
anova(mod3, mod1)

# same, but with syntax (not run)
specific3 <- "S1 = 7,9,10,11,13,15,17,18,21,22,24,27,31
             S2 = 1,3,6,8,16,29
             S3 = 2,4,5,12,14,19,20,23,25,26,28,30"
# mod3b <- bfactor(data, specific3)
# anova(mod3b)
```

```
#####
```

```
# mixed itemtype example
```

```
# simulate data
```

```
a <- matrix(c(
  1,0.5,NA,
  1,0.5,NA,
  1,0.5,NA,
  1,0.5,NA,
  1,0.5,NA,
  1,0.5,NA,
  1,0.5,NA,
  1,0.5,NA,
  1,NA,0.5,
  1,NA,0.5,
  1,NA,0.5,
  1,NA,0.5,
  1,NA,0.5,
  1,NA,0.5,
  1,NA,0.5),ncol=3,byrow=TRUE)
```

```
d <- matrix(c(
  -1.0,NA,NA,
  -1.5,NA,NA,
  1.5,NA,NA,
  0.0,NA,NA,
  2.5,1.0,-1,
  3.0,2.0,-0.5,
  3.0,2.0,-0.5,
  3.0,2.0,-0.5,
  2.5,1.0,-1,
  2.0,0.0,NA,
  -1.0,NA,NA,
  -1.5,NA,NA,
  1.5,NA,NA,
  0.0,NA,NA),ncol=3,byrow=TRUE)
```

```
items <- rep('2PL', 14)
items[5:10] <- 'graded'
```

```
sigma <- diag(3)
```

```

dataset <- simdata(a,d,5000,itemtype=items,sigma=sigma)
itemstats(dataset)

specific <- "S1 = 1-7
            S2 = 8-14"
simmod <- bfactor(dataset, specific)
coef(simmod, simplify=TRUE)

#####
# General testlet response model (Wainer, 2007)

# simulate data
set.seed(1234)
a <- matrix(0, 12, 4)
a[,1] <- rlnorm(12, .2, .3)
ind <- 1
for(i in 1:3){
  a[ind:(ind+3),i+1] <- a[ind:(ind+3),1]
  ind <- ind+4
}
print(a)
d <- rnorm(12, 0, .5)
sigma <- diag(c(1, .5, 1, .5))
dataset <- simdata(a,d,2000,itemtype=rep('2PL', 12),sigma=sigma)
itemstats(dataset)

# estimate by applying constraints and freeing the latent variances
specific <- "S1 = 1-4
            S2 = 5-8
            S3 = 9-12"
model <- "G = 1-12
CONSTRAIN = (1, a1, a2), (2, a1, a2), (3, a1, a2), (4, a1, a2),
            (5, a1, a3), (6, a1, a3), (7, a1, a3), (8, a1, a3),
            (9, a1, a4), (10, a1, a4), (11, a1, a4), (12, a1, a4)
COV = S1*S1, S2*S2, S3*S3"

simmod <- bfactor(dataset, specific, model)
coef(simmod, simplify=TRUE)

# Constrained testlet model (Bradlow, 1999)
model2 <- "G = 1-12
CONSTRAIN = (1, a1, a2), (2, a1, a2), (3, a1, a2), (4, a1, a2),
            (5, a1, a3), (6, a1, a3), (7, a1, a3), (8, a1, a3),
            (9, a1, a4), (10, a1, a4), (11, a1, a4), (12, a1, a4),
            (GROUP, COV_22, COV_33, COV_44)
COV = S1*S1, S2*S2, S3*S3"

simmod2 <- bfactor(dataset, specific, model2)
coef(simmod2, simplify=TRUE)
anova(simmod2, simmod)

```

```
#####
# Two-tier model

# simulate data
set.seed(1234)
a <- matrix(c(
  0,1,0.5,NA,NA,
  0,1,0.5,NA,NA,
  0,1,0.5,NA,NA,
  0,1,0.5,NA,NA,
  0,1,0.5,NA,NA,
  0,1,NA,0.5,NA,
  0,1,NA,0.5,NA,
  0,1,NA,0.5,NA,
  1,0,NA,0.5,NA,
  1,0,NA,0.5,NA,
  1,0,NA,0.5,NA,
  1,0,NA,NA,0.5,
  1,0,NA,NA,0.5,
  1,0,NA,NA,0.5,
  1,0,NA,NA,0.5,
  1,0,NA,NA,0.5),ncol=5,byrow=TRUE)

d <- matrix(rnorm(16))
items <- rep('2PL', 16)

sigma <- diag(5)
sigma[1,2] <- sigma[2,1] <- .4
dataset <- simdata(a,d,2000,itemtype=items,sigma=sigma)
itemstats(dataset)

specific <- "S1 = 1-5
            S2 = 6-11
            S3 = 12-16"

model <- '
  G1 = 1-8
  G2 = 9-16
  COV = G1*G2'

# quadpts dropped for faster estimation, but not as precise
simmod <- bfactor(dataset, specific, model, quadpts = 9, TOL = 1e-3)
coef(simmod, simplify=TRUE)
summary(simmod)
itemfit(simmod, QMC=TRUE)
M2(simmod, QMC=TRUE)
residuals(simmod, QMC=TRUE)

## End(Not run)
```

Bock1997

*Description of Bock 1997 data***Description**

A 3-item tabulated data set extracted from Table 3 in Chapter Two.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Bock, R. D. (1997). The Nominal Categories Model. In van der Linden, W. J. & Hambleton, R. K. *Handbook of modern item response theory*. New York: Springer.

Examples

```
## Not run:
dat <- expand.table(Bock1997)
head(dat)
itemstats(dat, use_ts=FALSE)

mod <- mirt(dat, 1, 'nominal')

# reproduce table 3 in Bock (1997)
fs <- round(fscores(mod, verbose = FALSE, full.scores = FALSE)[,c('F1', 'SE_F1')], 2)
fttd <- residuals(mod, type = 'exp')
table <- data.frame(fttd[, -ncol(fttd)], fs)
table

mod <- mirt(dat, 1, 'nominal')
coef(mod)

## End(Not run)
```

boot.LR

*Parametric bootstrap likelihood-ratio test***Description**

Given two fitted models, compute a parametric bootstrap test to determine whether the less restrictive models fits significantly better than the more restricted model. Note that this hypothesis test also works when prior parameter distributions are included for either model. Function can be run in parallel after using a suitable [mirtCluster](#) definition.

Usage

```
boot.LR(mod, mod2, R = 1000, verbose = TRUE)
```

Arguments

mod	an estimated model object, more constrained than mod2
mod2	an estimated model object
R	number of parametric bootstraps to use.
verbose	logical; include additional information in the console?

Value

a p-value evaluating whether the more restrictive model fits significantly worse than the less restrictive model

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
## Not run:

# standard
dat <- expand.table(LSAT7)
mod1 <- mirt(dat, 1)
mod2 <- mirt(dat, 1, '3PL')

# standard LR test
anova(mod1, mod2)

# bootstrap LR test (run in parallel to save time)
if(interactive()) mirtCluster()
boot.LR(mod1, mod2, R=200)

## End(Not run)
```

boot.mirt

*Calculate bootstrapped standard errors for estimated models***Description**

Given an internal mirt object estimate the bootstrapped standard errors. It may be beneficial to run the computations using multi-core architecture (e.g., the `parallel` package). Parameters are organized from the freely estimated values in `mod2values(x)` (equality constraints will also be returned in the bootstrapped estimates).

Usage

```
boot.mirt(x, R = 100, boot.fun = NULL, technical = NULL, ...)
```

Arguments

<code>x</code>	an estimated model object
<code>R</code>	number of draws to use (passed to the <code>boot()</code> function)
<code>boot.fun</code>	a user-defined function used to extract the information from the bootstrap fitted models. Must be of the form <code>boot.fun(x)</code> , where <code>x</code> is the bootstrap fitted model under investigation, and the return must be a numeric vector. If omitted a default function will be defined internally that returns the estimated parameters from the <code>mod</code> object, resulting in bootstrapped parameter estimate results
<code>technical</code>	technical arguments passed to estimation engine. See mirt for details
<code>...</code>	additional arguments to be passed on to <code>boot(...)</code> and <code>mirt</code> 's estimation engine

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```
## Not run:

# standard
mod <- mirt(Science, 1)
booted <- boot.mirt(mod, R=20)
plot(booted)
booted

#run in parallel using snow back-end using all available cores
mod <- mirt(Science, 1)
```

```

booted <- boot.mirt(mod, parallel = 'snow', ncpus = parallel::detectCores())
booted

####
# bootstrapped CIs for standardized factor loadings
boot.fun <- function(mod){
  so <- summary(mod, verbose=FALSE)
  as.vector(so$rotF)
}

# test to see if it works before running
boot.fun(mod)

# run
booted.loads <- boot.mirt(mod, boot.fun=boot.fun)
booted.loads

## End(Not run)

```

coef-method

Extract raw coefs from model object

Description

Return a list (or data.frame) of raw item and group level coefficients. Note that while the output to the console is rounded to three digits, the returned list of objects is not. Hence, elements from `cfs` `<- coef(mod)`; `cfs[[1]]` will contain the non-rounded results (useful for simulations).

Usage

```

## S4 method for signature 'SingleGroupClass'
coef(
  object,
  CI = 0.95,
  printSE = FALSE,
  rotate = "none",
  Target = NULL,
  IRTpars = FALSE,
  rawug = FALSE,
  as.data.frame = FALSE,
  simplify = FALSE,
  unique = FALSE,
  verbose = TRUE,
  ...
)

```

Arguments

object	an object of class <code>SingleGroupClass</code> , <code>MultipleGroupClass</code> , or <code>MixedClass</code>
CI	the amount of converged used to compute confidence intervals; default is 95 percent confidence intervals
printSE	logical; print the standard errors instead of the confidence intervals? When <code>IRTpars = TRUE</code> then the delta method will be used to compute the associated standard errors from mirt's default slope-intercept form
rotate	see summary method for details. The default rotation is 'none'
Target	a dummy variable matrix indicting a target rotation pattern
IRTpars	logical; convert slope intercept parameters into traditional IRT parameters? Only applicable to unidimensional models or models with simple structure (i.e., only one non-zero slope). If a suitable ACOV estimate was computed in the fitted model, and <code>printSE = FALSE</code> , then suitable CIs will be included based on the delta method (where applicable)
rawug	logical; return the untransformed internal g and u parameters? If <code>FALSE</code> , g and u's are converted with the original format along with delta standard errors
as.data.frame	logical; convert list output to a <code>data.frame</code> instead?
simplify	logical; if all items have the same parameter names (indicating they are of the same class) then they are collapsed to a matrix, and a list of length 2 is returned containing a matrix of item parameters and group-level estimates
unique	return the vector of uniquely estimated parameters
verbose	logical; allow information to be printed to the console?
...	additional arguments to be passed

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[summary-method](#)

Examples

```
## Not run:
dat <- expand.table(LSAT7)
x <- mirt(dat, 1)
coef(x)
coef(x, IRTpars = TRUE)
coef(x, simplify = TRUE)

#with computed information matrix
x <- mirt(dat, 1, SE = TRUE)
coef(x)
coef(x, printSE = TRUE)
```

```

coef(x, as.data.frame = TRUE)

#two factors
x2 <- mirt(Science, 2)
coef(x2)
coef(x2, rotate = 'varimax')

## End(Not run)

```

createGroup

Create a user defined group-level object with correct generic functions

Description

Initializes the proper S4 class and methods necessary for mirt functions to use in estimation for defining customized group-level functions. To use the defined objects pass to the `mirt(..., customGroup = OBJECT)` command, and ensure that the class parameters are properly labelled.

Usage

```

createGroup(
  par,
  est,
  den,
  nfact,
  standardize = FALSE,
  gr = NULL,
  hss = NULL,
  gen = NULL,
  lbound = NULL,
  ubound = NULL,
  derivType = "Richardson"
)

```

Arguments

<code>par</code>	a named vector of the starting values for the parameters
<code>est</code>	a logical vector indicating which parameters should be freely estimated by default
<code>den</code>	the probability density function given the Theta/ability values. First input contains a vector of all the defined parameters and the second input must be a matrix called Theta. Function also must return a numeric vector object corresponding to the associated densities for each row in the Theta input
<code>nfact</code>	number of factors required for the model. E.g., for unidimensional models with only one dimension of integration <code>nfact = 1</code>

standardize	logical; use standardization of the quadrature table method proposed by Woods and Thissen (2006)? If TRUE, the logical elements named 'MEAN_1' and 'COV_11' can be included in the parameter vector, and when these values are set to FALSE in the est input the E-table will be standardized to these fixed values (e.g., <code>par <- c(a1=1, d=0, MEAN_1=0, COV_11=1)</code> with <code>est <- c(TRUE, TRUE, FALSE, FALSE)</code> will standardize the E-table to have a 0 mean and unit variance)
gr	gradient function (vector of first derivatives) of the log-likelihood used in estimation. The function must be of the form <code>gr(x, Theta)</code> , where <code>x</code> is the object defined by <code>createGroup()</code> and <code>Theta</code> is a matrix of latent trait parameters
hss	Hessian function (matrix of second derivatives) of the log-likelihood used in estimation. If not specified a numeric approximation will be used. The input is identical to the <code>gr</code> argument
gen	a function used when <code>GenRandomPars = TRUE</code> is passed to the estimation function to generate random starting values. Function must be of the form <code>function(object) ...</code> and must return a vector with properties equivalent to the <code>par</code> object. If NULL, parameters will remain at the defined starting values by default
lbound	optional vector indicating the lower bounds of the parameters. If not specified then the bounds will be set to -Inf
ubound	optional vector indicating the lower bounds of the parameters. If not specified then the bounds will be set to Inf
derivType	if the <code>gr</code> or <code>hss</code> terms are not specified this type will be used to obtain them numerically. Default is 'Richardson'

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
# normal density example, N(mu, sigma^2)
den <- function(obj, Theta) dnorm(Theta, obj@par[1], sqrt(obj@par[2]))
par <- c(mu = 0, sigma2 = .5)
est <- c(FALSE, TRUE)
lbound <- c(-Inf, 0)
grp <- createGroup(par, est, den, nfact = 1, lbound=lbound)

dat <- expand.table(LSAT6)
mod <- mirt(dat, 1, 'Rasch')
modcustom <- mirt(dat, 1, 'Rasch', customGroup=grp)

coef(mod)
coef(modcustom)
```

createItem

*Create a user defined item with correct generic functions***Description**

Initializes the proper S4 class and methods necessary for `mirt` functions to use in estimation. To use the defined objects pass to the `mirt(..., customItems = list())` command, and ensure that the classes are properly labelled and unique in the list. Additionally, the input `mirt(..., customItemsData = list())` can also be included to specify additional item-level information to better recycle custom-item definitions (e.g., for supplying varying Q-matrices), where the `list` input must have the same length as the number of items. For further examples regarding how this function can be used for fitting unfolding-type models see Liu and Chalmers (2018).

Usage

```
createItem(
  name,
  par,
  est,
  P,
  gr = NULL,
  hss = NULL,
  gen = NULL,
  lbound = NULL,
  ubound = NULL,
  derivType = "Richardson",
  derivType.hss = "Richardson",
  bytecompile = TRUE
)
```

Arguments

<code>name</code>	a character indicating the item class name to be defined
<code>par</code>	a named vector of the starting values for the parameters
<code>est</code>	a logical vector indicating which parameters should be freely estimated by default
<code>P</code>	the probability trace function for all categories (first column is category 1, second category two, etc). First input contains a vector of all the item parameters, the second input must be a matrix called <code>Theta</code> , the third input must be the number of categories called <code>ncat</code> , and (optionally) a fourth argument termed <code>itemdata</code> may be included containing further users specification information. The last optional input is to be utilized within the estimation functions such as <code>mirt</code> via the list input <code>customItemsData</code> to more naturally recycle custom-item definitions. Therefore, these inputs must be of the form <code>function(par, Theta, ncat){...}</code> or

	<pre>function(par, Theta, ncat, itemdata){...}</pre> to be valid; however, the names of the arguments is not relevant. Finally, this function must return a <code>matrix</code> object of category probabilities, where the columns represent each respective category
gr	gradient function (vector of first derivatives) of the log-likelihood used in estimation. The function must be of the form <code>gr(x, Theta)</code> , where <code>x</code> is the object defined by <code>createItem()</code> and <code>Theta</code> is a matrix of latent trait parameters. Tabulated (EM) or raw (MHRM) data are located in the <code>x@dat</code> slot, and are used to form the complete data log-likelihood. If not specified a numeric approximation will be used
hss	Hessian function (matrix of second derivatives) of the log-likelihood used in estimation. If not specified a numeric approximation will be used (required for the MH-RM algorithm only). The input is identical to the <code>gr</code> argument
gen	a function used when <code>GenRandomPars = TRUE</code> is passed to the estimation function to generate random starting values. Function must be of the form <code>function(object) ...</code> and must return a vector with properties equivalent to the <code>par</code> object. If <code>NULL</code> , parameters will remain at the defined starting values by default
lbound	optional vector indicating the lower bounds of the parameters. If not specified then the bounds will be set to <code>-Inf</code>
ubound	optional vector indicating the upper bounds of the parameters. If not specified then the bounds will be set to <code>Inf</code>
derivType	if the <code>gr</code> term is not specified this type will be used to obtain the gradient numerically or symbolically. Default is the 'Richardson' extrapolation method; see numerical_deriv for details and other options. If 'symbolic' is supplied then the gradient will be computed using a symbolical approach (potentially the most accurate method, though may fail depending on how the <code>P</code> function was defined)
derivType.hss	if the <code>hss</code> term is not specified this type will be used to obtain the Hessian numerically. Default is the 'Richardson' extrapolation method; see numerical_deriv for details and other options. If 'symbolic' is supplied then the Hessian will be computed using a symbolical approach (potentially the most accurate method, though may fail depending on how the <code>P</code> function was defined)
bytecompile	logical; where applicable, byte compile the functions provided? Default is <code>TRUE</code> to provide

Details

The `summary()` function will not return proper standardized loadings since the function is not sure how to handle them (no slopes could be defined at all!). Instead loadings of .001 are filled in as place-holders.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Liu, C.-W. and Chalmers, R. P. (2018). Fitting item response unfolding models to Likert-scale data using mirt in R. *PLoS ONE*, 13, 5. doi:10.1371/journal.pone.0196292

Examples

```
## Not run:

name <- 'old2PL'
par <- c(a = .5, b = -2)
est <- c(TRUE, TRUE)
P.old2PL <- function(par, Theta, ncat){
  a <- par[1]
  b <- par[2]
  P1 <- 1 / (1 + exp(-1*a*(Theta - b)))
  cbind(1-P1, P1)
}

x <- createItem(name, par=par, est=est, P=P.old2PL)

# So, let's estimate it!
dat <- expand.table(LSAT7)
sv <- mirt(dat, 1, c(rep('2PL',4), 'old2PL'), customItems=list(old2PL=x), pars = 'values')
tail(sv) #looks good
mod <- mirt(dat, 1, c(rep('2PL',4), 'old2PL'), customItems=list(old2PL=x))
coef(mod)
mod2 <- mirt(dat, 1, c(rep('2PL',4), 'old2PL'), customItems=list(old2PL=x), method = 'MHRM')
coef(mod2)

# same definition as above, but using symbolic derivative computations
# (can be more accurate/stable)
xs <- createItem(name, par=par, est=est, P=P.old2PL, derivType = 'symbolic')
mod <- mirt(dat, 1, c(rep('2PL',4), 'old2PL'), customItems=list(old2PL=xs))
coef(mod, simplify=TRUE)

# several secondary functions supported
M2(mod, calcNull=FALSE)
itemfit(mod)
fscores(mod, full.scores=FALSE)
plot(mod)

# fit the same model, but specify gradient function explicitly (use of a browser() may be helpful)
gr <- function(x, Theta){
  # browser()
  a <- x@par[1]
  b <- x@par[2]
  P <- probtrace(x, Theta)
  PQ <- apply(P, 1, prod)
  r_P <- x@dat / P
}
```

```

    grad <- numeric(2)
    grad[2] <- sum(-a * PQ * (r_P[,2] - r_P[,1]))
    grad[1] <- sum((Theta - b) * PQ * (r_P[,2] - r_P[,1]))

    ## check with internal numerical form to be safe
    # numerical_deriv(x@par[x@est], mirt:::EML, obj=x, Theta=Theta)
    grad
  }

x <- createItem(name, par=par, est=est, P=P.old2PL, gr=gr)
mod <- mirt(dat, 1, c(rep('2PL',4), 'old2PL'), customItems=list(old2PL=x))
coef(mod, simplify=TRUE)

### non-linear
name <- 'nonlin'
par <- c(a1 = .5, a2 = .1, d = 0)
est <- c(TRUE, TRUE, TRUE)
P.nonlin <- function(par,Theta, ncat=2){
  a1 <- par[1]
  a2 <- par[2]
  d <- par[3]
  P1 <- 1 / (1 + exp(-1*(a1*Theta + a2*Theta^2 + d)))
  cbind(1-P1, P1)
}

x2 <- createItem(name, par=par, est=est, P=P.nonlin)

mod <- mirt(dat, 1, c(rep('2PL',4), 'nonlin'), customItems=list(nonlin=x2))
coef(mod)

### nominal response model (Bock 1972 version)
Tnom.dev <- function(ncat) {
  T <- matrix(1/ncat, ncat, ncat - 1)
  diag(T[-1, ]) <- diag(T[-1, ]) - 1
  return(T)
}

name <- 'nom'
par <- c(alp=c(3,0,-3),gam=rep(.4,3))
est <- rep(TRUE, length(par))
P.nom <- function(par, Theta, ncat){
  alp <- par[1:(ncat-1)]
  gam <- par[ncat:length(par)]
  a <- Tnom.dev(ncat) %*% alp
  c <- Tnom.dev(ncat) %*% gam
  z <- matrix(0, nrow(Theta), ncat)
  for(i in 1:ncat)
    z[,i] <- a[i] * Theta + c[i]
  P <- exp(z) / rowSums(exp(z))
  P
}

nom1 <- createItem(name, par=par, est=est, P=P.nom)

```

```

nommod <- mirt(Science, 1, 'nom1', customItems=list(nom1=nom1))
coef(nommod)
Tnom.dev(4) %%% coef(nommod)[[1]][1:3] #a
Tnom.dev(4) %%% coef(nommod)[[1]][4:6] #d

## End(Not run)

```

deAyala

*Description of deAyala data***Description**

Mathematics data from de Ayala (2009; pg. 14); 5 item dataset in table format.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

de Ayala, R. J. (2009). *The theory and practice of item response theory*. Guilford Press.

Examples

```

## Not run:
dat <- expand.table(deAyala)
head(dat)
itemstats(dat)

## End(Not run)

```

DeltaMethod

*Numerical derivative version of delta method***Description**

Delta method using numerical derivatives (via [numerical_deriv](#)) for the provided function. Convenient when target transformation function is easier to automate programmatically instead of using explicit formula or math expressions. Can also be useful for checking analytic results.

Usage

```
DeltaMethod(fn, par, acov, ...)
```

Arguments

<code>fn</code>	a function specifying the type of transformation to make for each new parameter of interest. Must be of the form <code>fn(par, ...)</code> , or more simply <code>fn(par)</code> , and return a numeric vector with one element
<code>par</code>	numerical vector passed to <code>fn(par)</code> (typically vector of MLEs)
<code>acov</code>	numeric matrix for the ACOV of the MLEs
<code>...</code>	additional arguments passed to <code>fn</code>

Value

returns a list of the transformed parameters, ACOV, and SEs

Examples

```
# Slightly modified example from ?msm::deltamethod
# Multiple linear regression,  $E(y) = \alpha + \beta_1 x + \beta_2 g$ 
x <- 1:100
g <- rep(0:1, each=50)
y <- rnorm(100, 4*x, 5)
toy.lm <- lm(y ~ x + g)
estmean <- coef(toy.lm)
estvar <- vcov(toy.lm)

# Estimate of  $(1 / (b_0 + b_1))$  and  $(1 / (b_0 + b_1 + b_2))$ 
1 / (estmean[1] + estmean[2])
1 / (estmean[1] + estmean[2] + estmean[3])

## Approximate standard error (uncomment to check)
# msm::deltamethod (~ 1 / (x1 + x2), estmean, estvar)
# msm::deltamethod (~ 1 / (x1 + x2 + x3), estmean, estvar)

# with DeltaMethod
fn <- function(par) 1 / sum(par[1:2])
fn2 <- function(par) 1 / sum(par[1:3])
DeltaMethod(fn, estmean, estvar)$se
DeltaMethod(fn2, estmean, estvar)$se

# index argument for easier flexibility
fn <- function(par, index) 1 / sum(par[index])
DeltaMethod(fn, estmean, estvar, index=1:2)$se
DeltaMethod(fn, estmean, estvar, index=1:3)$se
```

Description

This function runs the Wald and likelihood-ratio approaches for testing differential item functioning (DIF) with two or more groups. This is primarily a convenience wrapper to the [multipleGroup](#) function for performing standard DIF procedures. Independent models can be estimated in parallel by defining a parallel object with [mirtCluster](#), which will help to decrease the run time. For best results, the baseline model should contain a set of 'anchor' items and have freely estimated hyper-parameters in the focal groups.

Usage

```
DIF(
  MGmodel,
  which.par,
  scheme = "add",
  items2test = 1:extract.mirt(MGmodel, "nitems"),
  groups2test = "all",
  seq_stat = "SABIC",
  Wald = FALSE,
  p.adjust = "none",
  pairwise = FALSE,
  return_models = FALSE,
  return_seq_model = FALSE,
  max_run = Inf,
  plotdif = FALSE,
  type = "trace",
  simplify = TRUE,
  verbose = TRUE,
  ...
)
```

Arguments

MGmodel	an object returned from multipleGroup to be used as the reference model
which.par	a character vector containing the parameter names which will be inspected for DIF
scheme	type of DIF analysis to perform, either by adding or dropping constraints across groups. These can be: 'add' parameters in which.par will be constrained each item one at a time for items that are specified in items2test. This is beneficial when examining DIF from a model with parameters freely estimated across groups, and when inspecting differences via the Wald test 'drop' parameters in which.par will be freely estimated for items that are specified in items2test. This is useful when supplying an overly restrictive model and attempting to detect DIF with a slightly less restrictive model 'add_sequential' sequentially loop over the items being tested, and at the end of the loop treat DIF tests that satisfy the seq_stat criteria as invariant. The loop is then re-run on the remaining invariant items to determine if

they are now displaying DIF in the less constrained model, and when no new invariant item is found the algorithm stops and returns the items that displayed DIF. Note that the DIF statistics are relative to this final, less constrained model which includes the DIF effects

'drop_sequential' sequentially loop over the items being tested, and at the end of the loop treat items that violate the `seq_stat` criteria as demonstrating DIF. The loop is then re-run, leaving the items that previously demonstrated DIF as variable across groups, and the remaining test items that previously showed invariance are re-tested. The algorithm stops when no more items showing DIF are found and returns the items that displayed DIF. Note that the DIF statistics are relative to this final, less constrained model which includes the DIF effects

<code>items2test</code>	a numeric vector, or character vector containing the item names, indicating which items will be tested for DIF. In models where anchor items are known, omit them from this vector. For example, if items 1 and 2 are anchors in a 10 item test, then <code>items2test = 3:10</code> would work for testing the remaining items (important to remember when using sequential schemes)
<code>groups2test</code>	a character vector indicating which groups to use in the DIF testing investigations. Default is <code>'all'</code> , which uses all group information to perform joint hypothesis tests of DIF (for a two group setup these result in pair-wise tests). For example, if the group names were <code>'g1'</code> , <code>'g2'</code> and <code>'g3'</code> , and DIF was only to be investigated between group <code>'g1'</code> and <code>'g3'</code> then pass <code>groups2test = c('g1', 'g3')</code>
<code>seq_stat</code>	select a statistic to test for in the sequential schemes. Potential values are (in descending order of power) <code>'AIC'</code> , <code>'SABIC'</code> , <code>'HQ'</code> , and <code>'BIC'</code> . If a numeric value is input that ranges between 0 and 1, the <code>'p'</code> value will be tested (e.g., <code>seq_stat = .05</code> will test for the difference of $p < .05$ in the add scheme, or $p > .05$ in the drop scheme), along with the specified <code>p.adjust</code> input
<code>Wald</code>	logical; perform Wald tests for DIF instead of likelihood ratio test?
<code>p.adjust</code>	string to be passed to the <code>p.adjust</code> function to adjust p-values. Adjustments are located in the <code>adj_p</code> element in the returned list
<code>pairwise</code>	logical; perform pairwise tests between groups when the number of groups is greater than 2? Useful as quickly specified post-hoc tests
<code>return_models</code>	logical; return estimated model objects for further analysis? Default is <code>FALSE</code>
<code>return_seq_model</code>	logical; on the last iteration of the sequential schemes, return the fitted multiple-group model containing the freely estimated parameters indicative of DIF? This is generally only useful when <code>scheme = 'add_sequential'</code> . Default is <code>FALSE</code>
<code>max_run</code>	a number indicating the maximum number of cycles to perform in sequential searches. The default is to perform search until no further DIF is found
<code>plotdif</code>	logical; create item plots for items that are displaying DIF according to the <code>seq_stat</code> criteria? Only available for <code>'add'</code> type schemes
<code>type</code>	the type of plot argument passed to <code>plot()</code> . Default is <code>'trace'</code> , though another good option is <code>'infotrace'</code> . For ease of viewing, the <code>facet_item</code> argument to <code>mirt's plot()</code> function is set to <code>TRUE</code>

simplify	logical; simplify the output by returning a data.frame object with the differences between AIC, BIC, etc, as well as the chi-squared test (X2) and associated df and p-values
verbose	logical print extra information to the console?
...	additional arguments to be passed to <code>multipleGroup</code> and <code>plot</code>

Details

Generally, the pre-computed baseline model should have been configured with two estimation properties: 1) a set of 'anchor' items, where the anchor items have various parameters that have been constrained to be equal across the groups, and 2) contain freely estimated latent mean and variance terms in all but one group (the so-called 'reference' group). These two properties help to fix the metric of the groups so that item parameter estimates do not contain latent distribution characteristics.

Value

a `mirt_df` object with the information-based criteria for DIF, though this may be changed to a list output when `return_models` or `simplify` are modified. As well, a silent 'DIF_coefficients' attribute is included to view the item parameter differences between the groups

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Chalmers, R. P., Counsell, A., and Flora, D. B. (2016). It might not make a big DIF: Improved Differential Test Functioning statistics that account for sampling variability. *Educational and Psychological Measurement*, 76, 114-140. doi:10.1177/0013164415584576

See Also

`multipleGroup`, `DRF`

Examples

```
## Not run:

# simulate data where group 2 has a smaller slopes and more extreme intercepts
set.seed(12345)
a1 <- a2 <- matrix(abs(rnorm(15,1,.3)), ncol=1)
d1 <- d2 <- matrix(rnorm(15,0,.7), ncol=1)
a2[1:2, ] <- a1[1:2, ]/3
d1[c(1,3), ] <- d2[c(1,3), ]/4
head(data.frame(a.group1 = a1, a.group2 = a2, d.group1 = d1, d.group2 = d2))
itemtype <- rep('2PL', nrow(a1))
N <- 1000
```

```

dataset1 <- simdata(a1, d1, N, itemtype)
dataset2 <- simdata(a2, d2, N, itemtype, mu = .1, sigma = matrix(1.5))
dat <- rbind(dataset1, dataset2)
group <- c(rep('D1', N), rep('D2', N))

#### no anchors, all items tested for DIF by adding item constraints one item at a time.
# define a parallel cluster (optional) to help speed up internal functions
if(interactive()) mirtCluster()

# Information matrix with Oakes' identity (not controlling for latent group differences)
# NOTE: Without properly equating the groups the following example code is not testing for DIF,
#       # but instead reflects a combination of DIF + latent-trait distribution effects
model <- multipleGroup(dat, 1, group, SE = TRUE)

# Likelihood-ratio test for DIF (as well as model information)
dif <- DIF(model, c('a1', 'd'))
dif

# function silently includes "DIF_coefficients" attribute to view
# the IRT parameters post-completion
extract.mirt(dif, "DIF_coefficients")

# same as above, but using Wald tests with Benjamini & Hochberg adjustment
DIF(model, c('a1', 'd'), Wald = TRUE, p.adjust = 'fdr')

# equate the groups by assuming the last 5 items have no DIF
itemnames <- colnames(dat)
model <- multipleGroup(dat, 1, group, SE = TRUE,
  invariance = c(itemnames[11:ncol(dat)], 'free_means', 'free_var'))

# test whether adding slopes and intercepts constraints results in DIF. Plot items showing DIF
resulta1d <- DIF(model, c('a1', 'd'), plotdif = TRUE, items2test=1:10)
resulta1d

# test whether adding only slope constraints results in DIF for all items
DIF(model, 'a1', items2test=1:10)

# Determine whether it's a1 or d parameter causing DIF (could be joint, however)
(a1s <- DIF(model, 'a1', items2test = 1:3))
(ds <- DIF(model, 'd', items2test = 1:3))

### drop down approach (freely estimating parameters across groups) when
### specifying a highly constrained model with estimated latent parameters
model_constrained <- multipleGroup(dat, 1, group,
  invariance = c(colnames(dat), 'free_means', 'free_var'))
dropdown <- DIF(model_constrained, c('a1', 'd'), scheme = 'drop')
dropdown

# View silent "DIF_coefficients" attribute
extract.mirt(dropdown, "DIF_coefficients")

### sequential schemes (add constraints)

```

```

### sequential searches using SABIC as the selection criteria
# starting from completely different models
stepup <- DIF(model, c('a1', 'd'), scheme = 'add_sequential',
              items2test=1:10)

stepup

# step down procedure for highly constrained model
stepdown <- DIF(model_constrained, c('a1', 'd'), scheme = 'drop_sequential')
stepdown

# view final MG model (only useful when scheme is 'add_sequential')
updated_mod <- DIF(model, c('a1', 'd'), scheme = 'add_sequential',
                  return_seq_model=TRUE)
plot(updated_mod, type='trace')

#####
# Multi-group example

a1 <- a2 <- a3 <- matrix(abs(rnorm(15,1,.3))), ncol=1)
d1 <- d2 <- d3 <- matrix(rnorm(15,0,.7),ncol=1)
a2[1:2, ] <- a1[1:2, ]/3
d3[c(1,3), ] <- d2[c(1,3), ]/4
head(data.frame(a.group1 = a1, a.group2 = a2, a.group3 = a3,
                d.group1 = d1, d.group2 = d2, d.group3 = d3))
itemtype <- rep('2PL', nrow(a1))
N <- 1000
dataset1 <- simdata(a1, d1, N, itemtype)
dataset2 <- simdata(a2, d2, N, itemtype, mu = .1, sigma = matrix(1.5))
dataset3 <- simdata(a3, d3, N, itemtype, mu = .2)
dat <- rbind(dataset1, dataset2, dataset3)
group <- gl(3, N, labels = c('g1', 'g2', 'g3'))

# equate the groups by assuming the last 5 items have no DIF
itemnames <- colnames(dat)
model <- multipleGroup(dat, group=group, SE=TRUE,
                      invariance = c(itemnames[11:ncol(dat)], 'free_means', 'free_var'))
coef(model, simplify=TRUE)

# omnibus tests
dif <- DIF(model, which.par = c('a1', 'd'), items2test=1:9)
dif

# pairwise post-hoc tests for items flagged via omnibus tests
dif.posthoc <- DIF(model, which.par = c('a1', 'd'), items2test=1:2,
                  pairwise = TRUE)
dif.posthoc

# further probing for df = 1 tests, this time with Wald tests
DIF(model, which.par = c('a1'), items2test=1:2, pairwise = TRUE,
    Wald=TRUE)
DIF(model, which.par = c('d'), items2test=1:2, pairwise = TRUE,
    Wald=TRUE)

```

```
## End(Not run)
```

```
DiscreteClass-class    Class "DiscreteClass"
```

Description

Defines the object returned from `mdirt`.

Slots

Call: function call

Data: list of data, sometimes in different forms

Options: list of estimation options

Fit: a list of fit information

Model: a list of model-based information

ParObjects: a list of the S4 objects used during estimation

OptimInfo: a list of arguments from the optimization process

Internals: a list of internal arguments for secondary computations (inspecting this object is generally not required)

vcov: a matrix represented the asymptotic covariance matrix of the parameter estimates

time: a data.frame indicating the breakdown of computation times in seconds

Methods

print signature(x = "DiscreteClass")

show signature(object = "DiscreteClass")

anova signature(object = "DiscreteClass")

coef signature(x = "DiscreteClass")

summary signature(object = "DiscreteClass")

residuals signature(object = "DiscreteClass")

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

draw_parameters	<i>Draw plausible parameter instantiations from a given model</i>
-----------------	---

Description

Draws plausible parameters from a model using parametric sampling (if the information matrix was computed) or via bootstrap sampling. Primarily for use with the [DRF](#) function.

Usage

```
draw_parameters(
  mod,
  draws,
  method = c("parametric", "bootstrap"),
  redraws = 20,
  verbose = FALSE,
  ...
)
```

Arguments

mod	estimated single or multiple-group model
draws	number of draws to obtain
method	type of plausible values to obtain. Can be 'parametric', for the parametric sampling scheme which uses the estimated information matrix, or 'bootstrap' to obtain values from the boot function. Default is 'parametric'
redraws	number of redraws to perform when the given parameteric sample does not satisfy the upper and lower parameter bounds. If a valid set cannot be found within this number of draws then an error will be thrown
verbose	logical; include additional information in the console?
...	additional arguments to be passed

Value

returns a draws x p matrix of plausible parameters, where each row corresponds to a single set

Examples

```
## Not run:
set.seed(1234)
n <- 40
N <- 500

# only first 5 items as anchors
model <- 'F = 1-40
          CONSTRAINTB = (1-5, a1), (1-5, d)'
```

```

a <- matrix(1, n)
d <- matrix(rnorm(n), n)
group <- c(rep('Group_1', N), rep('Group_2', N))

## -----
# groups completely equal
dat1 <- simdata(a, d, N, itemtype = 'dich')
dat2 <- simdata(a, d, N, itemtype = 'dich')
dat <- rbind(dat1, dat2)
mod <- multipleGroup(dat, model, group=group, SE=TRUE,
                     invariance=c('free_means', 'free_var'))

param_set <- draw_parameters(mod, 100)
head(param_set)

## End(Not run)

```

DRF

Differential Response Functioning statistics

Description

Function performs various omnibus differential item (DIF), bundle (DBF), and test (DTF) functioning procedures on an object estimated with `multipleGroup()`. The compensatory and non-compensatory statistics provided are described in Chalmers (2018), which generally can be interpreted as IRT generalizations of the SIBTEST and CSIBTEST statistics. For hypothesis tests, these measures require the ACOV matrix to be computed in the fitted multiple-group model (otherwise, sets of plausible draws from the posterior are explicitly required).

Usage

```

DRF(
  mod,
  draws = NULL,
  focal_items = 1L:extract.mirt(mod, "nitems"),
  param_set = NULL,
  den.type = "marginal",
  best_fitting = FALSE,
  CI = 0.95,
  npts = 1000,
  quadpts = NULL,
  theta_lim = c(-6, 6),
  Theta_nodes = NULL,
  plot = FALSE,
  DIF = FALSE,
  DIF.cats = FALSE,

```

```

groups2test = "all",
pairwise = FALSE,
simplify = TRUE,
p.adjust = "none",
par.strip.text = list(cex = 0.7),
par.settings = list(strip.background = list(col = "#9ECAE1"), strip.border = list(col =
  "black")),
auto.key = list(space = "right", points = FALSE, lines = TRUE),
verbose = TRUE,
...
)

```

Arguments

<code>mod</code>	a multipleGroup object which estimated only 2 groups
<code>draws</code>	a number indicating how many draws to take to form a suitable multiple imputation or bootstrap estimate of the expected test scores (100 or more). If <code>boot = FALSE</code> , requires an estimated parameter information matrix. Returns a list containing the bootstrap/imputation distribution and null hypothesis test for the sDRF statistics
<code>focal_items</code>	a character/numeric vector indicating which items to include in the DRF tests. The default uses all of the items (note that including anchors in the focal items has no effect because they are exactly equal across groups). Selecting fewer items will result in tests of 'differential bundle functioning'
<code>param_set</code>	an N x p matrix of parameter values drawn from the posterior (e.g., using the parametric sampling approach, bootstrap, or MCMC). If supplied, then these will be used to compute the DRF measures. Can be much more efficient to pre-compute these values if DIF, DBF, or DTF are being evaluated within the same model (especially when using the bootstrap method). See draw_parameters
<code>den.type</code>	character specifying how the density of the latent traits is computed. Default is 'marginal' to include the proportional information from both groups, 'focal' for just the focal group, and 'reference' for the reference group
<code>best_fitting</code>	logical; use the best fitting parametric distribution (Gaussian by default) that was used at the time of model estimation? This will result in much fast computations, however the results are more dependent upon the underlying modelling assumptions. Default is FALSE, which uses the empirical histogram approach
<code>CI</code>	range of confidence interval when using draws input
<code>npts</code>	number of points to use for plotting. Default is 1000
<code>quadpts</code>	number of quadrature nodes to use when constructing DRF statistics. Default is extracted from the input model object
<code>theta_lim</code>	lower and upper limits of the latent trait (theta) to be evaluated, and is used in conjunction with <code>quadpts</code> and <code>npts</code>
<code>Theta_nodes</code>	an optional matrix of Theta values to be evaluated in the draws for the sDRF statistics. However, these values are not averaged across, and instead give the bootstrap confidence intervals at the respective Theta nodes. Useful when following up a large sDRF or uDRF statistic, for example, to determine where the

	difference between the test curves are large (while still accounting for sampling variability). Returns a matrix with observed variability
plot	logical; plot the 'sDRF' functions for the evaluated sDBF or sDTF values across the integration grid or, if DIF = TRUE, the selected items as a faceted plot of individual items? If plausible parameter sets were obtained/supplied then imputed confidence intervals will be included
DIF	logical; return a list of item-level imputation properties using the DRF statistics? These can generally be used as a DIF detection method and as a graphical display for understanding DIF within each item
DIF.cats	logical; same as DIF = TRUE, however computations will be performed on each item category probability functions rather than the score functions. Only useful for understanding DIF in polytomous items
groups2test	when more than 2 groups are being investigated which two groups should be used in the effect size comparisons?
pairwise	logical; perform pairwise computations when the applying to multi-group settings
simplify	logical; attempt to simplify the output rather than returning larger lists?
p.adjust	string to be passed to the p.adjust function to adjust p-values. Adjustments are located in the <code>adj_pvals</code> element in the returned list. Only applicable when DIF = TRUE
par.strip.text	plotting argument passed to lattice
par.settings	plotting argument passed to lattice
auto.key	plotting argument passed to lattice
verbose	logical; include additional information in the console?
...	additional arguments to be passed to lattice

Details

The effect sizes estimates by the DRF function are

$$sDRF = \int [S(C|\Psi^{(R)}, \theta)S(C|\Psi^{(F)}, \theta)]f(\theta)d\theta,$$

$$uDRF = \int |S(C|\Psi^{(R)}, \theta)S(C|\Psi^{(F)}, \theta)|f(\theta)d\theta,$$

and

$$dDRF = \sqrt{\int [S(C|\Psi^{(R)}, \theta)S(C|\Psi^{(F)}, \theta)]^2 f(\theta)d\theta}$$

where $S(\cdot)$ are the scoring equations used to evaluate the model-implied difference between the focal and reference group. The $f(\theta)$ terms can either be estimated from the posterior via an empirical histogram approach (default), or can use the best fitting prior distribution that is obtain post-convergence (default is a Guassian distribution). Note that, in comparison to Chalmers (2018), the focal group is the leftmost scoring function while the reference group is the rightmost scoring function. This is largely to keep consistent with similar effect size statistics, such as SIBTEST, DFIT, Wainer's measures of impact, etc, which in general can be seen as special-case estimators of this family.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R. P. (2018). Model-Based Measures for Detecting and Quantifying Response Bias. *Psychometrika*, 83(3), 696-732. doi:[10.1007/s1133601896269](https://doi.org/10.1007/s1133601896269)

See Also

[multipleGroup](#), [DIF](#)

Examples

```
## Not run:

set.seed(1234)
n <- 30
N <- 500

# only first 5 items as anchors
model <- 'F = 1-30
          CONSTRAINB = (1-5, a1), (1-5, d)'

a <- matrix(1, n)
d <- matrix(rnorm(n), n)
group <- c(rep('Group_1', N), rep('Group_2', N))

## -----
# groups completely equal
dat1 <- simdata(a, d, N, itemtype = 'dich')
dat2 <- simdata(a, d, N, itemtype = 'dich')
dat <- rbind(dat1, dat2)
mod <- multipleGroup(dat, model, group=group, SE=TRUE,
                    invariance=c('free_means', 'free_var'))

plot(mod)
plot(mod, which.items = 6:10) #DBF
plot(mod, type = 'itemscore')
plot(mod, type = 'itemscore', which.items = 10:15)

# empirical histogram approach
DRF(mod)
DRF(mod, focal_items = 6:10) #DBF
DRF(mod, DIF=TRUE)
DRF(mod, DIF=TRUE, focal_items = 10:15)

# Best-fitting Gaussian distributions
DRF(mod, best_fitting=TRUE)
DRF(mod, focal_items = 6:10, best_fitting=TRUE) #DBF
DRF(mod, DIF=TRUE, best_fitting=TRUE)
DRF(mod, DIF=TRUE, focal_items = 10:15, best_fitting=TRUE)
```

```

DRF(mod, plot = TRUE)
DRF(mod, focal_items = 6:10, plot = TRUE) #DBF
DRF(mod, DIF=TRUE, plot = TRUE)
DRF(mod, DIF=TRUE, focal_items = 10:15, plot = TRUE)

if(interactive()) mirtCluster()
DRF(mod, draws = 500)
DRF(mod, draws = 500, best_fitting=TRUE)
DRF(mod, draws = 500, plot=TRUE)

# pre-draw parameter set to save computations
# (more useful when using non-parametric bootstrap)
param_set <- draw_parameters(mod, draws = 500)
DRF(mod, focal_items = 6, param_set=param_set) #DIF test
DRF(mod, DIF=TRUE, param_set=param_set) #DIF test
DRF(mod, focal_items = 6:10, param_set=param_set) #DBF test
DRF(mod, param_set=param_set) #DTF test

DRF(mod, focal_items = 6:10, draws=500) #DBF test
DRF(mod, focal_items = 10:15, draws=500) #DBF test

DIFs <- DRF(mod, draws = 500, DIF=TRUE)
print(DIFs)
DRF(mod, draws = 500, DIF=TRUE, plot=TRUE)

DIFs <- DRF(mod, draws = 500, DIF=TRUE, focal_items = 6:10)
print(DIFs)
DRF(mod, draws = 500, DIF=TRUE, focal_items = 6:10, plot = TRUE)

DRF(mod, DIF=TRUE, focal_items = 6)
DRF(mod, draws=500, DIF=TRUE, focal_items = 6)

# evaluate specific values for sDRF
Theta_nodes <- matrix(seq(-6,6,length.out = 100))

sDTF <- DRF(mod, Theta_nodes=Theta_nodes)
head(sDTF)
sDTF <- DRF(mod, Theta_nodes=Theta_nodes, draws=200)
head(sDTF)

# sDIF (isolate single item)
sDIF <- DRF(mod, Theta_nodes=Theta_nodes, focal_items=6)
head(sDIF)
sDIF <- DRF(mod, Theta_nodes=Theta_nodes, focal_items = 6, draws=200)
head(sDIF)

## -----
## random slopes and intercepts for 15 items, and latent mean difference
## (no systematic DTF should exist, but DIF will be present)
set.seed(1234)
dat1 <- simdata(a, d, N, itemtype = 'dich', mu=.50, sigma=matrix(1.5))
dat2 <- simdata(a + c(numeric(15), rnorm(n-15, 0, .25)),
               d + c(numeric(15), rnorm(n-15, 0, .5)), N, itemtype = 'dich')

```

```

dat <- rbind(dat1, dat2)
mod1 <- multipleGroup(dat, 1, group=group)
plot(mod1)
DRF(mod1) #does not account for group differences! Need anchors

mod2 <- multipleGroup(dat, model, group=group, SE=TRUE,
                      invariance=c('free_means', 'free_var'))
plot(mod2)

# significant DIF in multiple items...
# DIF(mod2, which.par=c('a1', 'd'), items2test=16:30)
DRF(mod2)
DRF(mod2, draws=500) #non-sig DTF due to item cancellation

## -----
## systematic differing slopes and intercepts (clear DTF)
set.seed(1234)
dat1 <- simdata(a, d, N, itemtype = 'dich', mu=.50, sigma=matrix(1.5))
dat2 <- simdata(a + c(numeric(15), rnorm(n-15, 1, .25)),
               d + c(numeric(15), rnorm(n-15, 1, .5)),
               N, itemtype = 'dich')
dat <- rbind(dat1, dat2)
mod3 <- multipleGroup(dat, model, group=group, SE=TRUE,
                      invariance=c('free_means', 'free_var'))
plot(mod3) #visable DTF happening

# DIF(mod3, c('a1', 'd'), items2test=16:30)
DRF(mod3) #unsigned bias. Signed bias (group 2 scores higher on average)
DRF(mod3, draws=500)
DRF(mod3, draws=500, plot=TRUE) #multiple DRF areas along Theta

# plot the DIF
DRF(mod3, draws=500, DIF=TRUE, plot=TRUE)

# evaluate specific values for sDRF
Theta_nodes <- matrix(seq(-6,6,length.out = 100))
sDTF <- DRF(mod3, Theta_nodes=Theta_nodes, draws=200)
head(sDTF)

# DIF
sDIF <- DRF(mod3, Theta_nodes=Theta_nodes, focal_items = 30, draws=200)
car::some(sDIF)

## -----
# polytomous example
# simulate data where group 2 has a different slopes/intercepts
set.seed(4321)
a1 <- a2 <- matrix(rlnorm(20,.2,.3))
a2[c(16:17, 19:20),] <- a1[c(16:17, 19:20),] + c(-.5, -.25, .25, .5)

# for the graded model, ensure that there is enough space between the intercepts,
# otherwise closer categories will not be selected often
diffs <- t(apply(matrix(runif(20*4, .3, 1), 20), 1, cumsum))

```

```

diffs <- -(diffs - rowMeans(diffs))
d1 <- d2 <- diffs + rnorm(20)
rownames(d1) <- rownames(d2) <- paste0('Item.', 1:20)
d2[16:20,] <- d1[16:20,] + matrix(c(-.5, -.5, -.5, -.5,
                                   1, 0, 0, -1,
                                   .5, .5, -.5, -.5,
                                   1, .5, 0, -1,
                                   .5, .5, .5, .5), byrow=TRUE, nrow=5)

tail(data.frame(a.group1 = a1, a.group2 = a2), 6)
list(d.group1 = d1[15:20,], d.group2 = d2[15:20,])

itemtype <- rep('graded', nrow(a1))
N <- 600
dataset1 <- simdata(a1, d1, N, itemtype)
dataset2 <- simdata(a2, d2, N, itemtype, mu = -.25, sigma = matrix(1.25))
dat <- rbind(dataset1, dataset2)
group <- c(rep('D1', N), rep('D2', N))

# item 1-10 as anchors
mod <- multipleGroup(dat, group=group, SE=TRUE,
                    invariance=c(colnames(dat)[1:10], 'free_means', 'free_var'))
coef(mod, simplify=TRUE)
plot(mod)
plot(mod, type='itemscore')

# DIF tests vis Wald method
DIF(mod, items2test=11:20,
    which.par=c('a1', paste0('d', 1:4)),
    Wald=TRUE, p.adjust='holm')

DRF(mod)
DRF(mod, DIF=TRUE, focal_items=11:20)
DRF(mod, DIF.cats=TRUE, focal_items=11:20)

## -----
### multidimensional DTF

set.seed(1234)
n <- 50
N <- 1000

# only first 5 items as anchors within each dimension
model <- 'F1 = 1-25
        F2 = 26-50
        COV = F1*F2
        CONSTRAINB = (1-5, a1), (1-5, 26-30, d), (26-30, a2)'

a <- matrix(c(rep(1, 25), numeric(50), rep(1, 25)), n)
d <- matrix(rnorm(n), n)
group <- c(rep('Group_1', N), rep('Group_2', N))
Cov <- matrix(c(1, .5, .5, 1.5), 2)
Mean <- c(0, 0.5)

```

```

# groups completely equal
dat1 <- simdata(a, d, N, itemtype = 'dich', sigma = cov2cor(Cov))
dat2 <- simdata(a, d, N, itemtype = 'dich', sigma = Cov, mu = Mean)
dat <- rbind(dat1, dat2)
mod <- multipleGroup(dat, model, group=group, SE=TRUE,
                     invariance=c('free_means', 'free_var'))
coef(mod, simplify=TRUE)
plot(mod, degrees = c(45,45))
DRF(mod)

# some intercepts slightly higher in Group 2
d2 <- d
d2[c(10:15, 31:35)] <- d2[c(10:15, 31:35)] + 1
dat1 <- simdata(a, d, N, itemtype = 'dich', sigma = cov2cor(Cov))
dat2 <- simdata(a, d2, N, itemtype = 'dich', sigma = Cov, mu = Mean)
dat <- rbind(dat1, dat2)
mod <- multipleGroup(dat, model, group=group, SE=TRUE,
                     invariance=c('free_means', 'free_var'))
coef(mod, simplify=TRUE)
plot(mod, degrees = c(45,45))

DRF(mod)
DRF(mod, draws = 500)

## End(Not run)

```

Description

Function performs various omnibus differential test functioning procedures on an object estimated with `multipleGroup()`. If the latent means/covariances are suspected to differ then the input object should contain a set of 'anchor' items to ensure that only differential test features are being detected rather than group differences. Returns signed (average area above and below) and unsigned (total area) statistics, with descriptives such as the percent average bias between group total scores for each statistic. If a grid of Theta values is passed, these can be evaluated as well to determine specific DTF location effects. For best results, the baseline model should contain a set of 'anchor' items and have freely estimated hyper-parameters in the focal groups. See [DIF](#) for details.

Usage

```

DTF(
  mod,
  draws = NULL,
  CI = 0.95,
  npts = 1000,

```

```

theta_lim = c(-6, 6),
Theta_nodes = NULL,
plot = "none",
auto.key = list(space = "right", points = FALSE, lines = TRUE),
...
)

```

Arguments

<code>mod</code>	a multipleGroup object which estimated only 2 groups
<code>draws</code>	a number indicating how many draws to take to form a suitable multiple imputation estimate of the expected test scores (usually 100 or more). Returns a list containing the imputation distribution and null hypothesis test for the sDTF statistic
<code>CI</code>	range of confidence interval when using draws input
<code>npts</code>	number of points to use in the integration. Default is 1000
<code>theta_lim</code>	lower and upper limits of the latent trait (theta) to be evaluated, and is used in conjunction with <code>npts</code>
<code>Theta_nodes</code>	an optional matrix of Theta values to be evaluated in the draws for the sDTF statistic. However, these values are not averaged across, and instead give the bootstrap confidence intervals at the respective Theta nodes. Useful when following up a large uDTF/sDTF statistic to determine where the difference between the test curves are large (while still accounting for sampling variability). Returns a matrix with observed variability
<code>plot</code>	a character vector indicating which plot to draw. Possible values are 'none', 'func' for the test score functions, and 'sDTF' for the evaluated sDTF values across the integration grid. Each plot is drawn with imputed confidence envelopes
<code>auto.key</code>	logical; automatically generate key in lattice plot?
<code>...</code>	additional arguments to be passed to <code>lattice</code> and <code>boot</code>

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Chalmers, R. P., Counsell, A., and Flora, D. B. (2016). It might not make a big DIF: Improved Differential Test Functioning statistics that account for sampling variability. *Educational and Psychological Measurement*, 76, 114-140. doi:10.1177/0013164415584576

See Also

[multipleGroup](#), [DIF](#)

Examples

```
## Not run:
set.seed(1234)
n <- 30
N <- 500

# only first 5 items as anchors
model <- 'F = 1-30
        CONSTRAINB = (1-5, a1), (1-5, d)'

a <- matrix(1, n)
d <- matrix(rnorm(n), n)
group <- c(rep('Group_1', N), rep('Group_2', N))

## -----
# groups completely equal
dat1 <- simdata(a, d, N, itemtype = '2PL')
dat2 <- simdata(a, d, N, itemtype = '2PL')
dat <- rbind(dat1, dat2)
mod <- multipleGroup(dat, model, group=group, SE=TRUE,
                    invariance=c('free_means', 'free_var'))
plot(mod)

DTF(mod)
if(interactive()) mirtCluster()
DTF(mod, draws = 1000) #95% C.I. for sDTF containing 0. uDTF is very small
DTF(mod, draws = 1000, plot='sDTF') #sDTF 95% C.I.'s across Theta always include 0

## -----
## random slopes and intercepts for 15 items, and latent mean difference
## (no systematic DTF should exist, but DIF will be present)
set.seed(1234)
dat1 <- simdata(a, d, N, itemtype = '2PL', mu=.50, sigma=matrix(1.5))
dat2 <- simdata(a + c(numeric(15), runif(n-15, -.2, .2)),
               d + c(numeric(15), runif(n-15, -.5, .5)), N, itemtype = '2PL')
dat <- rbind(dat1, dat2)
mod1 <- multipleGroup(dat, 1, group=group)
plot(mod1) #does not account for group differences! Need anchors

mod2 <- multipleGroup(dat, model, group=group, SE=TRUE,
                    invariance=c('free_means', 'free_var'))
plot(mod2)

# significant DIF in multiple items....
# DIF(mod2, which.par=c('a1', 'd'), items2test=16:30)
DTF(mod2)
DTF(mod2, draws=1000) #non-sig DTF due to item cancellation

## -----
## systematic differing slopes and intercepts (clear DTF)
dat1 <- simdata(a, d, N, itemtype = '2PL', mu=.50, sigma=matrix(1.5))
dat2 <- simdata(a + c(numeric(15), rnorm(n-15, 1, .25)), d + c(numeric(15), rnorm(n-15, 1, .5)),
```

```

                                N, itemtype = '2PL')
dat <- rbind(dat1, dat2)
mod3 <- multipleGroup(dat, model, group=group, SE=TRUE,
                      invariance=c('free_means', 'free_var'))
plot(mod3) #visable DTF happening

# DIF(mod3, c('a1', 'd'), items2test=16:30)
DTF(mod3) #unsigned bias. Signed bias indicates group 2 scores generally higher on average
DTF(mod3, draws=1000)
DTF(mod3, draws=1000, plot='func')
DTF(mod3, draws=1000, plot='sDTF') #multiple DTF areas along Theta

# evaluate specific values for sDTF
Theta_nodes <- matrix(seq(-6,6,length.out = 100))
sDTF <- DTF(mod3, Theta_nodes=Theta_nodes)
head(sDTF)
sDTF <- DTF(mod3, Theta_nodes=Theta_nodes, draws=100)
head(sDTF)

## End(Not run)

```

empirical_ES

Empirical effect sizes based on latent trait estimates

Description

Computes effect size measures of differential item functioning and differential test/bundle functioning based on expected scores from Meade (2010). Item parameters from both reference and focal group are used in conjunction with focal group empirical theta estimates (and an assumed normally distributed theta) to compute expected scores.

Usage

```

empirical_ES(
  mod,
  Theta.focal = NULL,
  focal_items = 1L:extract.mirt(mod, "nitems"),
  DIF = TRUE,
  npts = 61,
  theta_lim = c(-6, 6),
  plot = FALSE,
  type = "b",
  par.strip.text = list(cex = 0.7),
  par.settings = list(strip.background = list(col = "#9ECAE1"), strip.border = list(col =
    "black")),
  ...
)

```

Arguments

<code>mod</code>	a <code>multipleGroup</code> object which estimated only 2 groups. The first group in this object is assumed to be the reference group by default (i.e., <code>ref.group = 1</code>), which conforms to the invariance arguments in multipleGroup
<code>Theta.focal</code>	an optional matrix of Theta values from the focal group to be evaluated. If not supplied the default values to fscores will be used in conjunction with the ... arguments passed
<code>focal_items</code>	a numeric vector indicating which items to include the tests. The default uses all of the items. Selecting fewer items will result in tests of 'differential bundle functioning' when <code>DIF = FALSE</code>
<code>DIF</code>	logical; return a data.frame of item-level imputation properties? If <code>FALSE</code> , only DBF and DTF statistics will be reported
<code>npts</code>	number of points to use in the integration. Default is 61
<code>theta_lim</code>	lower and upper limits of the latent trait (theta) to be evaluated, and is used in conjunction with <code>npts</code>
<code>plot</code>	logical; plot expected scores of items/test where expected scores are computed using focal group thetas and both focal and reference group item parameters
<code>type</code>	type of objects to draw in lattice; default plots both points and lines
<code>par.strip.text</code>	plotting argument passed to lattice
<code>par.settings</code>	plotting argument passed to lattice
<code>...</code>	additional arguments to be passed to fscores and xyplot

DIF

The default `DIF = TRUE` produces several effect sizes indices at the item level. Signed indices allow DIF favoring the focal group at one point on the theta distribution to cancel DIF favoring the reference group at another point on the theta distribution. Unsigned indices take the absolute value before summing or averaging, thus not allowing cancellation of DIF across theta.

SIDS Signed Item Difference in the Sample. The average difference in expected scores across the focal sample using both focal and reference group item parameters.

UIDS Unsigned Item Difference in the Sample. Same as SIDS except absolute value of expected scores is taken prior to averaging across the sample.

D-Max The maximum difference in expected scores in the sample.

ESSD Expected Score Standardized Difference. Cohen's D for difference in expected scores.

SIDN Signed Item Difference in a Normal distribution. Identical to SIDS but averaged across a normal distribution rather than the sample.

UIDN Unsigned Item Difference in a Normal distribution. Identical to UIDS but averaged across a normal distribution rather than the sample.

DBF/DTF

DIF = FALSE produces a series of test/bundle-level indices that are based on item-level indices.

STDS Signed Test Differences in the Sample. The sum of the SIDS across items.

UTDS Unsigned Test Differences in the Sample. The sum of the UIDS across items.

Stark's DTFR Stark's version of STDS using a normal distribution rather than sample estimated thetas.

UDTFR Unsigned Expected Test Scores Differences in the Sample. The difference in observed summed scale scores expected, on average, across a hypothetical focal group with a normally distributed theta, had DF been uniform in nature for all items

UETSDS Unsigned Expected Test Score Differences in the Sample. The hypothetical difference expected scale scores that would have been present if scale-level DF had been uniform across respondents (i.e., always favoring the focal group).

UETSDN Identical to UETSDS but computed using a normal distribution.

Test D-Max Maximum expected test score differences in the sample.

ETSSD Expected Test Score Standardized Difference. Cohen's D for expected test scores.

Author(s)

Adam Meade, with contributions by Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Meade, A. W. (2010). A taxonomy of effect size measures for the differential functioning of items and scales. *Journal of Applied Psychology*, 95, 728-743.

Examples

```
## Not run:

# no DIF
set.seed(12345)
a <- matrix(abs(rnorm(15,1,.3)), ncol=1)
d <- matrix(rnorm(15,0,.7),ncol=1)
itemtype <- rep('2PL', nrow(a))
N <- 1000
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a, d, N, itemtype, mu = .1, sigma = matrix(1.5))
dat <- rbind(dataset1, dataset2)

# ensure 'Ref' is the first group (and therefore reference group during estimation)
group <- factor(c(rep('Ref', N), rep('Focal', N)), levels = c('Ref', 'Focal'))

mod <- multipleGroup(dat, 1, group = group,
  invariance = c(colnames(dat)[1:5], 'free_means', 'free_var'))
coef(mod, simplify=TRUE)
```

```

empirical_ES(mod)
empirical_ES(mod, DIF=FALSE)
empirical_ES(mod, DIF=FALSE, focal_items = 10:15)

empirical_ES(mod, plot=TRUE)
empirical_ES(mod, plot=TRUE, DIF=FALSE)

###-----
# DIF
set.seed(12345)
a1 <- a2 <- matrix(abs(rnorm(15,1,.3)), ncol=1)
d1 <- d2 <- matrix(rnorm(15,0,.7),ncol=1)
a2[10:15,] <- a2[10:15,] + rnorm(6, 0, .3)
d2[10:15,] <- d2[10:15,] + rnorm(6, 0, .3)
itemtype <- rep('dich', nrow(a1))
N <- 1000
dataset1 <- simdata(a1, d1, N, itemtype)
dataset2 <- simdata(a2, d2, N, itemtype, mu = .1, sigma = matrix(1.5))
dat <- rbind(dataset1, dataset2)
group <- factor(c(rep('Ref', N), rep('Focal', N)), levels = c('Ref', 'Focal'))

mod <- multipleGroup(dat, 1, group = group,
  invariance = c(colnames(dat)[1:5], 'free_means', 'free_var'))
coef(mod, simplify=TRUE)

empirical_ES(mod)
empirical_ES(mod, DIF = FALSE)
empirical_ES(mod, plot=TRUE)
empirical_ES(mod, plot=TRUE, DIF=FALSE)

## End(Not run)

```

empirical_plot

Function to generate empirical unidimensional item and test plots

Description

Given a dataset containing item responses this function will construct empirical graphics using the observed responses to each item conditioned on the total score. When individual item plots are requested then the total score will be formed without the item of interest (i.e., the total score without that item).

Usage

```

empirical_plot(
  data,
  which.items = NULL,
  type = "prop",

```

```

smooth = FALSE,
formula = resp ~ s(TS, k = 5),
main = NULL,
par.strip.text = list(cex = 0.7),
par.settings = list(strip.background = list(col = "#9ECAE1"), strip.border = list(col =
  "black")),
auto.key = list(space = "right", points = FALSE, lines = TRUE),
...
)

```

Arguments

<code>data</code>	a data.frame or matrix of item responses (see mirt for typical input)
<code>which.items</code>	a numeric vector indicating which items to plot in a faceted image plot. If NULL then empirical test plots will be constructed instead
<code>type</code>	character vector specifying type of plot to draw. When <code>which.item</code> is NULL can be 'prop' (default) or 'hist', otherwise can be 'prop' (default) or 'boxplot'
<code>smooth</code>	logical; include a GAM smoother instead of the raw proportions? Default is FALSE
<code>formula</code>	formula used for the GAM smoother
<code>main</code>	the main title for the plot. If NULL an internal default will be used
<code>par.strip.text</code>	plotting argument passed to lattice
<code>par.settings</code>	plotting argument passed to lattice
<code>auto.key</code>	plotting argument passed to lattice
<code>...</code>	additional arguments to be passed to lattice and <code>coef()</code>

Details

Note that these types of plots should only be used for unidimensional tests with monotonically increasing item response functions. If monotonicity is not true for all items, however, then these plots may serve as a visual diagnostic tool so long as the majority of items are indeed monotonic.

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[itemstats](#), [itemplot](#), [itemGAM](#)

Examples

```

## Not run:

SAT12[SAT12 == 8] <- NA
data <- key2binary(SAT12,

```

```

key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))

# test plot
empirical_plot(data)
empirical_plot(data, type = 'hist')
empirical_plot(data, type = 'hist', breaks=20)

# items 1, 2 and 5
empirical_plot(data, c(1, 2, 5))
empirical_plot(data, c(1, 2, 5), smooth = TRUE)
empirical_plot(data, c(1, 2, 5), type = 'boxplot')

# replace weird looking items with unscored versions for diagnostics
empirical_plot(data, 32)
data[,32] <- SAT12[,32]
empirical_plot(data, 32)
empirical_plot(data, 32, smooth = TRUE)

## End(Not run)

```

empirical_rxx

Function to calculate the empirical (marginal) reliability

Description

Given secondary latent trait estimates and their associated standard errors returned from [fcores](#), compute the empirical reliability.

Usage

```
empirical_rxx(Theta_SE, T_as_X = FALSE)
```

Arguments

Theta_SE	a matrix of latent trait estimates returned from fcores with the options <code>full.scores = TRUE</code> and <code>full.scores.SE = TRUE</code>
T_as_X	logical; should the observed variance be equal to $\text{var}(X) = \text{var}(T) + E(E^2)$ or $\text{var}(X) = \text{var}(T)$ when computing empirical reliability estimates? Default (FALSE) uses the former

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[fscores](#), [marginal_rxx](#)

Examples

```
## Not run:

dat <- expand.table(deAyala)
itemstats(dat)
mod <- mirt(dat)

theta_se <- fscores(mod, full.scores.SE = TRUE)
empirical_rxx(theta_se)

theta_se <- fscores(mod, full.scores.SE = TRUE, method = 'ML')
empirical_rxx(theta_se)
empirical_rxx(theta_se, T_as_X = TRUE)

## End(Not run)
```

estfun.AllModelClass *Extract Empirical Estimating Functions*

Description

A function for extracting the empirical estimating functions of a fitted [mirt](#), [multipleGroup](#), [bfactor](#), or [mdirt](#) model. This is the derivative of the log-likelihood with respect to the parameter vector, evaluated at the observed (case-wise) data. In other words, this function returns the case-wise scores, evaluated at the fitted model parameters. Currently, models fitted via the EM or BL method are supported. For the computations, the internal Theta grid of the model is being used which was already used during the estimation of the model itself along with its matching normalized density.

Usage

```
estfun.AllModelClass(
  x,
  weights = extract.mirt(x, "survey.weights"),
  centering = FALSE
)
```

Arguments

x a fitted model object of class `SingleGroupClass`, `MultipleGroupClass`, or `DiscreteClass`

weights	by default, the survey.weights which were (optionally) specified when fitting the model are included to calculate the scores. If specified by the user, this should be a numeric vector of length equal to the total sample size. Note that if not all cases were weighted equally when fitting the model, the weights must be corrected by taking their square root if the scores are being used to compute the outer product of gradients (OPG) estimate of the variance-covariance matrix (see examples below).
centering	a boolean variable that allows the centering of the case-wise scores (i.e., setting their expected values to 0). If the case-wise scores were obtained from maximum likelihood estimates, this setting does not affect the result.

Value

An $n \times k$ matrix corresponding to n observations and k parameters

Author(s)

Lennart Schneider <lennart.sch@web.de> and Phil Chalmers; centering argument contributed by Rudolf Debelak (<rudolf.debelak@psychologie.uzh.ch>)

See Also

[mirt](#), [multipleGroup](#), [bfactor](#), [mdirt](#)

Examples

```
## Not run:
# fit a 2PL on the LSAT7 data and get the scores
mod1 <- mirt(expand.table(LSAT7), 1, SE = TRUE, SE.type = "crossprod")
sc1 <- estfun.AllModelClass(mod1)
# get the gradient
colSums(sc1)
# calculate the OPG estimate of the variance-covariance matrix "by hand"
vc1 <- vcov(mod1)
all.equal(crossprod(sc1), chol2inv(chol(vc1)), check.attributes = FALSE)

# Discrete group
modd <- mdirt(expand.table(LSAT7), 2, SE = TRUE, SE.type = "crossprod")
sc1 <- estfun.AllModelClass(modd)
# get the gradient
colSums(sc1)
# calculate the OPG estimate of the variance-covariance matrix "by hand"
vc1 <- vcov(modd)
all.equal(crossprod(sc1), chol2inv(chol(vc1)), check.attributes = FALSE)

# fit a multiple group 2PL and do the same as above
group <- rep(c("G1", "G2"), 500)
mod2 <- multipleGroup(expand.table(LSAT7), 1, group, SE = TRUE,
  SE.type = "crossprod")
sc2 <- estfun.AllModelClass(mod2)
colSums(sc2)
```

```

vc2 <- vcov(mod2)
all.equal(crossprod(sc2), chol2inv(chol(vc2)), check.attributes = FALSE)

# fit a bifactor model with 2 specific factors and do the same as above
mod3 <- bfactor(expand.table(LSAT7), c(2, 2, 1, 1, 2), SE = TRUE,
  SE.type = "crossprod")
sc3 <- estfun.AllModelClass(mod3)
colSums(sc3)
vc3 <- vcov(mod3)
all.equal(crossprod(sc3), chol2inv(chol(vc3)), check.attributes = FALSE)

# fit a 2PL not weighting all cases equally
survey.weights <- c(rep(2, sum(LSAT7$freq) / 2), rep(1, sum(LSAT7$freq) / 2))
survey.weights <- survey.weights / sum(survey.weights) * sum(LSAT7$freq)
mod4 <- mirt(expand.table(LSAT7), 1, SE = TRUE, SE.type = "crossprod",
  survey.weights = survey.weights)
sc4 <- estfun.AllModelClass(mod4,
  weights = extract.mirt(mod4, "survey.weights"))
# get the gradient
colSums(sc4)
# to calculate the OPG estimate of the variance-covariance matrix "by hand",
# the weights must be adjusted by taking their square root
sc4_crp <- estfun.AllModelClass(mod4,
  weights = sqrt(extract.mirt(mod4, "survey.weights")))
vc4 <- vcov(mod4)
all.equal(crossprod(sc4_crp), chol2inv(chol(vc4)), check.attributes = FALSE)

## End(Not run)

```

expand.table

Expand summary table of patterns and frequencies

Description

The `expand.table` function expands a summary table of unique response patterns to a full sized data-set. By default the response frequencies are assumed to be on rightmost column of the input data, though this can be modified.

Usage

```
expand.table(tabdata, freq = colnames(tabdata)[ncol(tabdata)], sample = FALSE)
```

Arguments

tabdata	An object of class <code>data.frame</code> or <code>matrix</code> with the unique response patterns and the number of frequencies in the rightmost column (though see <code>freq</code> for details on how to omit this column)
---------	---

freq	either a character vector specifying the column in tabdata to be used as the frequency count indicator for each response pattern (defaults to the right-most column) or a integer vector of length nrow(tabdata) specifying the frequency counts. When using the latter approach the tabdata input should not include any information regarding the counts, and instead should only include the unique response patterns themselves
sample	logical; randomly switch the rows in the expanded table? This does not change the expanded data, only the row locations

Value

Returns a numeric matrix with all the response patterns.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
data(LSAT7)
head(LSAT7) # frequency in right-most column
LSAT7full <- expand.table(LSAT7)
head(LSAT7full)
dim(LSAT7full)

# randomly switch rows in the expanded response table
LSAT7samp <- expand.table(LSAT7, sample = TRUE)
head(LSAT7samp)
colMeans(LSAT7full)
colMeans(LSAT7samp) #equal

#-----

## Not run:
# Generate data from separate response pattern matrix and freq vector
# The following uses Table 2.1 from de Ayala (2009)
f <- c(691,2280,242,235,158,184,1685,1053,134,462,92,65,571,79,87,41,1682,702,
      370,63,626,412,166,52,28,15,2095,1219,500,187,40,3385)

pat <- matrix(c(
  0, 0, 0, 0, 0,
  1, 0, 0, 0, 0,
  0, 1, 0, 0, 0,
  0, 0, 1, 0, 0,
  0, 0, 0, 1, 0,
  0, 0, 0, 0, 1,
```

```

1, 1, 0, 0, 0,
1, 0, 1, 0, 0,
0, 1, 1, 0, 0,
1, 0, 0, 1, 0,
0, 1, 0, 1, 0,
0, 0, 1, 1, 0,
1, 0, 0, 0, 1,
0, 1, 0, 0, 1,
0, 0, 1, 0, 1,
0, 0, 0, 1, 1,
1, 1, 1, 0, 0,
1, 1, 0, 1, 0,
1, 0, 1, 1, 0,
0, 1, 1, 1, 0,
1, 1, 0, 0, 1,
1, 0, 1, 0, 1,
1, 0, 0, 1, 1,
0, 1, 1, 0, 1,
0, 1, 0, 1, 1,
0, 0, 1, 1, 1,
1, 1, 1, 1, 0,
1, 1, 1, 0, 1,
1, 1, 0, 1, 1,
1, 0, 1, 1, 1,
0, 1, 1, 1, 1,
1, 1, 1, 1, 1), ncol=5, byrow=TRUE)

colnames(pat) <- paste0('Item.', 1:5)
head(pat)

table2.1 <- expand.table(pat, freq = f)
dim(table2.1)

## End(Not run)

```

expected.item

Function to calculate expected value of item

Description

Given an internal mirt object extracted from an estimated model compute the expected value for an item given the ability parameter(s).

Usage

```
expected.item(x, Theta, min = 0, include.var = FALSE)
```

Arguments

x	an extracted internal mirt object containing item information (see extract.item)
Theta	a vector (unidimensional) or matrix (multidimensional) of latent trait values
min	a constant value added to the expected values indicating the lowest theoretical category. Default is 0
include.var	logical; include the model-implied variance of the expected scores as well? When TRUE will return a list containing the expected values (E) and variances (VAR)

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

See Also

[extract.item](#), [expected.test](#)

Examples

```
mod <- mirt(Science, 1)
extr.2 <- extract.item(mod, 2)
Theta <- matrix(seq(-6,6, length.out=200))
expected <- expected.item(extr.2, Theta, min(Science[,1])) #min() of first item
head(data.frame(expected, Theta=Theta))

expected.item(extr.2, Theta, min(Science[,1]), include.var=TRUE)
```

expected.test

Function to calculate expected test score

Description

Given an estimated model compute the expected test score. Returns the expected values in the same form as the data used to estimate the model.

Usage

```

expected.test(
  x,
  Theta,
  group = NULL,
  mins = TRUE,
  individual = FALSE,
  which.items = NULL,
  probs.only = FALSE
)

```

Arguments

x	an estimated mirt object
Theta	a matrix of latent trait values (if a vector is supplied, will be coerced to a matrix with one column)
group	a number or character signifying which group the item should be extracted from (applies to 'MultipleGroupClass' objects only)
mins	logical; include the minimum value constants in the dataset. If FALSE, the expected values for each item are determined from the scoring 0:(ncat-1)
individual	logical; return tracelines for individual items?
which.items	an integer vector indicating which items to include in the expected test score. Default uses all possible items
probs.only	logical; return the probability for each category instead of traceline score functions? Only useful when individual=TRUE

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[expected.item](#)

Examples

```

## Not run:
dat <- expand.table(deAyala)
model <- 'F = 1-5
          CONSTRAIN = (1-5, a1)'
mod <- mirt(dat, model)

Theta <- matrix(seq(-6,6,.01))
tscore <- expected.test(mod, Theta)
tail(cbind(Theta, tscore))

# use only first two items (i.e., a bundle)

```

```
bscore <- expected.test(mod, Theta, which.items = 1:2)
tail(cbind(Theta, bscore))

# more low-level output (score and probability elements)
expected.test(mod, Theta, individual=TRUE)
expected.test(mod, Theta, individual=TRUE, probs.only=TRUE)

## End(Not run)
```

extract.group

Extract a group from a multiple group mirt object

Description

Extract a single group from an object defined by [multipleGroup](#).

Usage

```
extract.group(x, group)
```

Arguments

x	mirt model of class 'MultipleGroupClass'
group	the name of the group to extract

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

See Also

[extract.item](#), [extract.mirt](#)

Examples

```
## Not run:
set.seed(12345)
a <- matrix(abs(rnorm(15,1,.3))), ncol=1)
d <- matrix(rnorm(15,0,.7),ncol=1)
itemtype <- rep('2PL', nrow(a))
N <- 1000
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a, d, N, itemtype, mu = .1, sigma = matrix(1.5))
```

```

dat <- rbind(dataset1, dataset2)
group <- c(rep('D1', N), rep('D2', N))
models <- 'F1 = 1-15'

mod_configural <- multipleGroup(dat, models, group = group)
group.1 <- extract.group(mod_configural, 'D1') #extract first group
summary(group.1)
plot(group.1)

## End(Not run)

```

extract.item

Extract an item object from mirt objects

Description

Extract the internal mirt objects from any estimated model.

Usage

```
extract.item(x, item, group = NULL, drop.zeros = FALSE)
```

Arguments

x	mirt model of class 'SingleGroupClass' or 'MultipleGroupClass'
item	a number or character signifying which item to extract
group	a number signifying which group the item should be extracted from (applies to 'MultipleGroupClass' only)
drop.zeros	logical; drop slope values that are numerically close to zero to reduce dimensionality? Useful in objects returned from bfactor or other confirmatory models that contain several zero slopes

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

See Also

[extract.group](#), [extract.mirt](#)

Examples

```

## Not run:
mod <- mirt(Science, 1)
extr.1 <- extract.item(mod, 1)

## End(Not run)

```

extract.mirt	<i>Extract various elements from estimated model objects</i>
--------------	--

Description

A generic function to extract the internal objects from estimated models.

Usage

```
extract.mirt(x, what, item = 1)
```

Arguments

x	mirt model of class 'SingleGroupClass', 'MultipleGroupClass', 'MixedClass' or 'DiscreteGroupClass'
what	a string indicating what to extract
item	if necessary, which item to extract information from. Defaults to 1 if not specified

Details

Objects which can be extracted from mirt objects include:

logLik observed log-likelihood

logPrior log term contributed by prior parameter distributions

G2 goodness of fit statistic

df degrees of freedom

p p-value for G2 statistic

RMSEA root mean-square error of approximation based on G2

CFI CFI fit statistic

TLI TLI fit statistic

AIC AIC

BIC BIC

SABIC sample size adjusted BIC

HQ HQ

LLhistory EM log-likelihood history

tabdata a tabular version of the raw response data input. Frequencies are stored in freq

freq frequencies associated with tabdata

K an integer vector indicating the number of unique elements for each item

mins an integer vector indicating the lowest category found in the input data

model input model syntax

method estimation method used

itemtype a vector of item types for each respective item (e.g., 'graded', '2PL', etc)

itemnames a vector of item names from the input data

factorNames a vector of factor names from the model definition

rowID an integer vector indicating all valid row numbers used in the model estimation (when all cases are used this will be 1:nrow(data))

data raw input data of item responses

covdata raw input data of data used as covariates

tabdataalong similar to tabdata, however the responses have been transformed into dummy coded variables

fulldataalong analogous to tabdatafull, but for the raw input data instead of the tabulated frequencies

EMhistory if saved, extract the EM iteration history

exp_resp expected probability of the unique response patterns

survey.weights if supplied, the vector of survey weights used during estimation (NULL if missing)

converged a logical value indicating whether the model terminated within the convergence criteria

iterations number of iterations it took to reach the convergence criteria

nest number of freely estimated parameters

parvec vector containing uniquely estimated parameters

vcov parameter covariance matrix (associated with parvec)

condnum the condition number of the Hessian (if computed). Otherwise NA

constrain a list of item parameter constraints to indicate which item parameters were equal during estimation

Prior prior density distribution for the latent traits

thetaPosterior posterior distribution for latent traits when using EM algorithm

key if supplied, the data scoring key

nfact number of latent traits/factors

nitems number of items

ngroups number of groups

groupNames character vector of unique group names

group a character vector indicating the group membership

invariance a character vector indicating invariance input from [multipleGroup](#)

secondordertest a logical indicating whether the model passed the second-order test based on the Hessian matrix. Indicates whether model is a potential local maximum solution

SEMconv logical; check whether the supplemented EM information matrix converged. Will be NA if not applicable

time estimation time, broken into different sections

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[extract.group](#), [extract.item](#), [mod2values](#)

Examples

```
## Not run:
mod <- mirt(Science, 1)

extract.mirt(mod, 'logLik')
extract.mirt(mod, 'K') # unique categories for each item

#multiple group model
grp <- rep(c('G1', 'G2'), each = nrow(Science)/2)
mod2 <- multipleGroup(Science, 1, grp)

grp1 <- extract.group(mod2, 1) #extract single group model
extract.mirt(mod2, 'parvec')
extract.mirt(grp1, 'parvec')

## End(Not run)
```

fixedCalib

Fixed-item calibration method

Description

Implements the set of fixed-item calibration methods described by Kim (2006). The initial calibrated model must be fitted via [mirt](#), is currently limited to unidimensional models only, and should only be utilized when the new set of responses are obtained from a population with similar distributional characteristics in the latent traits. For more flexible calibration of items, including a fixed-item calibration variant involving anchor items for equating, see [multipleGroup](#).

Usage

```
fixedCalib(
  data,
  model = 1,
  old_mod,
```

```

    PAU = "MWU",
    NEMC = "MEM",
    technical = list(),
    ...
  )

```

Arguments

<code>data</code>	new data to be used for calibration. Note that to be consistent with the <code>mod</code> object, observed responses/NA placeholders must be included to link the item names used in the original <code>mod</code> definition (i.e., <code>extract.mirt(mod, what = 'itemnames')</code>)
<code>model</code>	type of model to fit for the complete dataset (not that for the fixed items in <code>old_mod</code> the factor loadings/constraints specified by the potential <code>mirt.model</code> specification is not relevant)
<code>old_mod</code>	a model of class <code>SingleGroupClass</code> fitted using <code>mirt</code>
<code>PAU</code>	prior ability update (PAU) approach. Supports none ("NWU"), one ("OWU"), and many ("MWU")
<code>NEMC</code>	number of EM cycles (NEMC) to use for the to-be-estimated parameters. Supports one ("OEM") and many ("MEM")
<code>technical</code>	list of technical estimation arguments (see <code>mirt</code> for details)
<code>...</code>	additional arguments to pass to <code>mirt</code>

References

Kim, S. (2006). A comparative study of IRT fixed parameter calibration methods. *Journal of Educational Measurement*, 4(43), 355-381.

See Also

`mirt`, `multipleGroup`

Examples

```

## Not run:

# single factor
set.seed(12345)
J <- 50
a <- matrix(abs(rnorm(J,1,.3)), ncol=1)
d <- matrix(rnorm(J,0,.7),ncol=1)
itemtype <- rep('2PL', nrow(a))

# calibration data theta ~ N(0,1)
N <- 3000
dataset1 <- simdata(a, d, N = N, itemtype=itemtype)

# new data (again, theta ~ N(0,1))
dataset2 <- simdata(a, d, N = 1000, itemtype=itemtype)

```

```

# last 40% of experimental items not given to calibration group
# (unobserved; hence removed)
dataset1 <- dataset1[,-c(J:(J*.6))]
head(dataset1)

#-----

# calibrated model from dataset1 only
mod <- mirt(dataset1, model = 1)
coef(mod, simplify=TRUE)

# No Prior Weights Updating and One EM Cycle (NWU-OEM)
NWU_OEM <- fixedCalib(dataset2, model=1, old_mod=mod, PAU='NWU', NEMC='OEM')
coef(NWU_OEM, simplify=TRUE)
data.frame(coef(NWU_OEM, simplify=TRUE)$items[,c('a1','d')],
           pop_a1=a, pop_d=d)
plot(NWU_OEM, type = 'empiricalhist')

# No Prior Weights Updating and Multiple EM Cycles (NWU-MEM)
NWU_MEM <- fixedCalib(dataset2, model = 1, old_mod = mod, PAU = 'NWU')
coef(NWU_MEM, simplify=TRUE)
data.frame(coef(NWU_MEM, simplify=TRUE)$items[,c('a1','d')],
           pop_a1=a, pop_d=d)
plot(NWU_MEM, type = 'empiricalhist')

# One Prior Weights Updating and One EM Cycle (OWU-OEM)
OWU_OEM <- fixedCalib(dataset2, model=1, old_mod=mod, PAU='OWU', NEMC="OEM")
coef(OWU_OEM, simplify=TRUE)
data.frame(coef(OWU_OEM, simplify=TRUE)$items[,c('a1','d')], pop_a1=a, pop_d=d)
plot(OWU_OEM, type = 'empiricalhist')

# One Prior Weights Updating and Multiple EM Cycles (OWU-MEM)
OWU_MEM <- fixedCalib(dataset2, model = 1, old_mod = mod, PAU = 'OWU')
coef(OWU_MEM, simplify=TRUE)
data.frame(coef(OWU_MEM, simplify=TRUE)$items[,c('a1','d')],
           pop_a1=a, pop_d=d)
plot(OWU_MEM, type = 'empiricalhist')

# Multiple Prior Weights Updating and Multiple EM Cycles (MWU-MEM)
MWU_MEM <- fixedCalib(dataset2, model = 1, old_mod = mod)
coef(MWU_MEM, simplify=TRUE)
data.frame(coef(MWU_MEM, simplify=TRUE)$items[,c('a1','d')],
           pop_a1=a, pop_d=d)
plot(MWU_MEM, type = 'empiricalhist')

# factor scores distribution check
fs <- fscores(MWU_MEM)
hist(fs)
c(mean_calib=mean(fs[1:N, ]), sd_calib=sd(fs[1:N, ]))
c(mean_exper=mean(fs[-c(1:N), ]), sd_exper=sd(fs[-c(1:N), ]))

#####

```

```
## Item length constraint example for each participant in the experimental
## items group. In this example, all participants were forced to have a test
## length of J=30, though the item pool had J=50 total items.

# new experimental data (relatively extreme, theta ~ N(.5,1.5))
dataset2 <- simdata(a, d, N = 1000, itemtype=itemtype,
  mu=.5, sigma=matrix(1.5))

# Add missing values to each participant in new dataset where individuals
# were randomly administered 10 experimental items, subject to the constraint
# that each participant received a test with J=30 items.
dataset2 <- t(apply(dataset2, 1, function(x){
  NA_precalib <- sample(1:30, 10)
  NA_experimental <- sample(31:50, 10)
  x[c(NA_precalib, NA_experimental)] <- NA
  x
}))
head(dataset2)

# check that all individuals had 30 items
all(rowSums(!is.na(dataset2)) == 30)

# Multiple Prior Weights Updating and Multiple EM Cycles (MWU-MEM)
MWU_MEM <- fixedCalib(dataset2, model = 1, old_mod = mod)
coef(MWU_MEM, simplify=TRUE)
data.frame(coef(MWU_MEM, simplify=TRUE)$items[,c('a1','d')],
  pop_a1=a, pop_d=d)
plot(MWU_MEM, type = 'empiricalhist')

## factor scores check
fs <- fscores(MWU_MEM)
hist(fs)
c(mean_calib=mean(fs[1:N, ]), sd_calib=sd(fs[1:N, ]))

## shrinkage, but generally different from calibrated sample
c(mean_exper=mean(fs[-c(1:N), ]), sd_exper=sd(fs[-c(1:N), ]))

## End(Not run)
```

fixef

Compute latent regression fixed effect expected values

Description

Create expected values for fixed effects parameters in latent regression models.

Usage

```
fixef(x)
```

Arguments

`x` an estimated model object from the `mixedmirt` or `mirt` function

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Chalmers, R. P. (2015). Extended Mixed-Effects Item Response Models with the MH-RM Algorithm. *Journal of Educational Measurement*, 52, 200-222. doi:10.1111/jedm.12072

See Also

`mirt`, `mixedmirt`

Examples

```
## Not run:

#simulate data
set.seed(1234)
N <- 1000

# covariates
X1 <- rnorm(N); X2 <- rnorm(N)
covdata <- data.frame(X1, X2)
Theta <- matrix(0.5 * X1 + -1 * X2 + rnorm(N, sd = 0.5))

#items and response data
a <- matrix(1, 20); d <- matrix(rnorm(20))
dat <- simdata(a, d, 1000, itemtype = '2PL', Theta=Theta)

#conditional model using X1 and X2 as predictors of Theta
mod1 <- mirt(dat, 1, 'Rasch', covdata=covdata, formula = ~ X1 + X2)

#latent regression fixed effects (i.e., expected values)
fe <- fixef(mod1)
head(fe)

# with mixedmirt()
mod1b <- mixedmirt(dat, covdata, 1, lr.fixed = ~ X1 + X2, fixed = ~ 0 + items)
fe2 <- fixef(mod1b)
head(fe2)

## End(Not run)
```

fscores	<i>Compute factor score estimates (a.k.a, ability estimates, latent trait estimates, etc)</i>
---------	---

Description

Computes MAP, EAP, ML (Embretson & Reise, 2000), EAP for sum-scores (Thissen et al., 1995), or WLE (Warm, 1989) factor scores with a multivariate normal prior distribution using equally spaced quadrature. EAP scores for models with more than three factors are generally not recommended since the integration grid becomes very large, resulting in slower estimation and less precision if the quadpts are too low. Therefore, MAP scores should be used instead of EAP scores for higher dimensional models. Multiple imputation variants are possible for each estimator if a parameter information matrix was computed, which are useful if the sample size/number of items were small. As well, if the model contained latent regression predictors this information will be used in computing MAP and EAP estimates (for these models, `full.scores=TRUE` will always be used). Finally, plausible value imputation is also available, and will also account for latent regression predictor effects.

Usage

```
fscores(
  object,
  method = "EAP",
  full.scores = TRUE,
  rotate = "oblimin",
  Target = NULL,
  response.pattern = NULL,
  append_response.pattern = FALSE,
  na.rm = FALSE,
  plausible.draws = 0,
  plausible.type = "normal",
  quadpts = NULL,
  item_weights = rep(1, extract.mirt(object, "nitems")),
  returnER = FALSE,
  T_as_X = FALSE,
  EAPsum.scores = FALSE,
  return.acov = FALSE,
  mean = NULL,
  cov = NULL,
  covdata = NULL,
  verbose = TRUE,
  full.scores.SE = FALSE,
  theta_lim = c(-6, 6),
  MI = 0,
  use_dentype_estimate = FALSE,
  QMC = FALSE,
  custom_den = NULL,
```

```

    custom_theta = NULL,
    min_expected = 1,
    max_theta = 20,
    start = NULL,
    ...
)

```

Arguments

object	a computed model object of class <code>SingleGroupClass</code> , <code>MultipleGroupClass</code> , or <code>DiscreteClass</code>
method	type of factor score estimation method. Can be: • "EAP" for the expected a-posteriori (default). For models fit using <code>mdirt</code> this will return the posterior classification probabilities • "MAP" for the maximum a-posteriori (i.e, Bayes modal) • "ML" for maximum likelihood • "WLE" for weighted likelihood estimation • "EAPsum" for the expected a-posteriori for each sum score • "plausible" for a single plausible value imputation for each case. This is equivalent to setting <code>plausible.draws = 1</code> • "classify" for the posteriori classification probabilities (only applicable when the input model was of class <code>MixtureClass</code>)
full.scores	if FALSE then a summary table with factor scores for each unique pattern is displayed as a formatted matrix object. Otherwise, a matrix of factor scores for each response pattern in the data is returned (default)
rotate	prior rotation to be used when estimating the factor scores. See summary-method for details. If the object is not an exploratory model then this argument is ignored
Target	target rotation; see summary-method for details
response.pattern	an optional argument used to calculate the factor scores and standard errors for a given response vector or matrix/data.frame
append_response.pattern	logical; should the inputs from <code>response.pattern</code> also be appended to the factor score output?
na.rm	logical; remove rows with any missing values? This is generally not required due to the nature of computing factors scores, however for the "EAPsum" method this may be necessary to ensure that the sum-scores correspond to the same composite score
plausible.draws	number of plausible values to draw for future researchers to perform secondary analyses of the latent trait scores. Typically used in conjunction with latent regression predictors (see mirt for details), but can also be generated when no predictor variables were modelled. If <code>plausible.draws</code> is greater than 0 a list of plausible values will be returned

<code>plausible.type</code>	type of plausible values to obtain. Can be either 'normal' (default) to use a normal approximation based on the ACOV matrix, or 'MH' to obtain Metropolis-Hastings samples from the posterior (silently passes object to <code>mirt</code> , therefore arguments like <code>technical</code> can be supplied to increase the number of burn-in draws and discarded samples)
<code>quadpts</code>	number of quadrature to use per dimension. If not specified, a suitable one will be created which decreases as the number of dimensions increases (and therefore for estimates such as EAP, will be less accurate). This is determined from the switch statement <code>quadpts <- switch(as.character(nfact), '1'=121, '2'=61, '3'=31, '4'=19, '5'=11, '6'=7, 5)</code>
<code>item_weights</code>	a user-defined weight vector used in the likelihood expressions to add more/less weight for a given observed response. Default is a vector of 1's, indicating that all the items receive the same weight
<code>returnER</code>	logical; return empirical reliability (also known as marginal reliability) estimates as a numeric values?
<code>T_as_X</code>	logical; should the observed variance be equal to $\text{var}(X) = \text{var}(T) + E(E^2)$ or $\text{var}(X) = \text{var}(T)$ when computing empirical reliability estimates? Default (FALSE) uses the former
<code>EAPsum.scores</code>	logical; include the model-implied expected values and variance for the item and total scores when using method = 'EAPsum' with <code>full.scores=FALSE</code> ? This information is included in the hidden 'fit' attribute which can be extracted via <code>attr(, 'fit')</code> for later use
<code>return.acov</code>	logical; return a list containing covariance matrices instead of factors scores? <code>impute = TRUE</code> not supported with this option
<code>mean</code>	a vector for custom latent variable means. If NULL, the default for 'group' values from the computed mirt object will be used
<code>cov</code>	a custom matrix of the latent variable covariance matrix. If NULL, the default for 'group' values from the computed mirt object will be used
<code>covdata</code>	when latent regression model has been fitted, and the <code>response.pattern</code> input is used to score individuals, then this argument is used to include the latent regression covariate terms for each row vector supplied to <code>response.pattern</code>
<code>verbose</code>	logical; print verbose output messages?
<code>full.scores.SE</code>	logical; when <code>full.scores == TRUE</code> , also return the standard errors associated with each respondent? Default is FALSE
<code>theta_lim</code>	lower and upper range to evaluate latent trait integral for each dimension. If omitted, a range will be generated automatically based on the number of dimensions
<code>MI</code>	a number indicating how many multiple imputation draws to perform. Default is 0, indicating that no MI draws will be performed
<code>use_dentype_estimate</code>	logical; if the density of the latent trait was estimated in the model (e.g., via Davidian curves or empirical histograms), should this information be used to compute the latent trait estimates? Only applicable for EAP-based estimates (EAP, EAPsum, and plausible)

QMC	logical; use quasi-Monte Carlo integration? If quadpts is omitted the default number of nodes is 5000
custom_den	a function used to define the integration density (if required). The NULL default assumes that the multivariate normal distribution with the 'GroupPars' hyper-parameters are used. At the minimum must be of the form: function(Theta, ...) where Theta is a matrix of latent trait values (will be a grid of values if method == 'EAPsum' or method == 'EAP', otherwise Theta will have only 1 row). Additional arguments may included and are caught through the fcores(...) input. The function <i>must</i> return a numeric vector of density weights (one for each row in Theta)
custom_theta	a matrix of custom integration nodes to use instead of the default, where each column corresponds to the respective dimension in the model
min_expected	when computing goodness of fit tests when method = 'EAPsum', this value is used to collapse across the conditioned total scores until the expected values are greater than this value. Note that this only affect the goodness of fit tests and not the returned EAP for sum scores table
max_theta	the maximum/minimum value any given factor score estimate will achieve using any modal estimator method (e.g., MAP, WLE, ML)
start	a matrix of starting values to use for iterative estimation methods. Default will start at a vector of 0's for each response pattern, or will start at the EAP estimates (unidimensional models only). Must be in the form that matches full.scores = FALSE (mostly used in the mirtCAT package)
...	additional arguments to be passed to nlm

Details

The function will return either a table with the computed scores and standard errors, the original data matrix with scores appended to the rightmost column, or the scores only. By default the latent means and covariances are determined from the estimated object, though these can be overwritten. Iterative estimation methods can be estimated in parallel to decrease estimation times if a [mirtCluster](#) object is available.

If the input object is a discrete latent class object estimated from [mdirt](#) then the returned results will be with respect to the posterior classification for each individual. The method inputs for 'DiscreteClass' objects may only be 'EAP', for posterior classification of each response pattern, or 'EAPsum' for posterior classification based on the raw sum-score. For more information on these algorithms refer to the [mirtCAT](#) package and the associated JSS paper (Chalmers, 2016).

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Chalmers, R. P. (2016). Generating Adaptive and Non-Adaptive Test Interfaces for Multidimensional Item Response Theory Applications. *Journal of Statistical Software*, 71(5), 1-39. doi:10.18637/jss.v071.i05

Embretson, S. E. & Reise, S. P. (2000). Item Response Theory for Psychologists. Erlbaum.

Thissen, D., Pommerich, M., Billeaud, K., & Williams, V. S. L. (1995). Item Response Theory for Scores on Tests Including Polytomous Items with Ordered Responses. *Applied Psychological Measurement*, 19, 39-49.

Warm, T. A. (1989). Weighted likelihood estimation of ability in item response theory. *Psychometrika*, 54, 427-450.

See Also

[averageMI](#)

Examples

```
mod <- mirt(Science)
tabscores <- fscores(mod, full.scores = FALSE)
head(tabscores)

# convert scores into expected total score information with 95% CIs
E.total <- expected.test(mod, Theta=tabscores[, 'F1'])
E.total_2.5 <- expected.test(mod, Theta=tabscores[, 'F1'] +
                             tabscores[, 'SE_F1'] * qnorm(.05/2))
E.total_97.5 <- expected.test(mod, Theta=tabscores[, 'F1'] +
                             tabscores[, 'SE_F1'] * qnorm(1-.05/2))

data.frame(Total_score=rowSums(tabscores[, 1:4]),
           E.total, E.total_2.5, E.total_97.5) |> head()

## Not run:
fullscores <- fscores(mod)
fullscores_with_SE <- fscores(mod, full.scores.SE=TRUE)
head(fullscores)
head(fullscores_with_SE)

# convert scores into expected total score information with 95% CIs
E.total <- expected.test(mod, Theta=fullscores[, 'F1'])
E.total_2.5 <- expected.test(mod, Theta=fullscores_with_SE[, 'F1'] +
                             fullscores_with_SE[, 'SE_F1'] * qnorm(.05/2))
E.total_97.5 <- expected.test(mod, Theta=fullscores_with_SE[, 'F1'] +
                             fullscores_with_SE[, 'SE_F1'] * qnorm(1-.05/2))

data.frame(Total_score=rowSums(Science),
           E.total, E.total_2.5, E.total_97.5) |> head()

# change method argument to use MAP estimates
fullscores <- fscores(mod, method='MAP')
head(fullscores)

# calculate MAP for a given response vector
```

```

fscores(mod, method='MAP', response.pattern = c(1,2,3,4))
# or matrix
fscores(mod, method='MAP', response.pattern = rbind(c(1,2,3,4), c(2,2,1,3)))

# return only the scores and their SEs
fscores(mod, method='MAP', response.pattern = c(1,2,3,4))

# use custom latent variable properties (diffuse prior for MAP is very close to ML)
fscores(mod, method='MAP', cov = matrix(1000), full.scores = FALSE)
fscores(mod, method='ML', full.scores = FALSE)

# EAPsum table of values based on total scores
(fs <- fscores(mod, method = 'EAPsum', full.scores = FALSE))

# convert expected counts back into marginal probability distribution
within(fs,
  `P(y)` <- expected / sum(observed))

# list of error VCOV matrices for EAPsum (works for other estimators as well)
acovs <- fscores(mod, method = 'EAPsum', full.scores = FALSE, return.acov = TRUE)
acovs

# WLE estimation, run in parallel using available cores
if(interactive()) mirtCluster()
head(fscores(mod, method='WLE', full.scores = FALSE))

# multiple imputation using 30 draws for EAP scores. Requires information matrix
mod <- mirt(Science, 1, SE=TRUE)
fs <- fscores(mod, MI = 30)
head(fs)

# plausible values for future work
pv <- fscores(mod, plausible.draws = 5)
lapply(pv, function(x) c(mean=mean(x), var=var(x), min=min(x), max=max(x)))

## define a custom_den function (*must* return a numeric vector).
# EAP with a uniform prior between -3 and 3
fun <- function(Theta, ...) as.numeric(dunif(Theta, min = -3, max = 3))
head(fscores(mod, custom_den = fun))

# compare EAP estimators with same modified prior
fun <- function(Theta, ...) as.numeric(dnorm(Theta, mean=.5))
head(fscores(mod, custom_den = fun))
head(fscores(mod, method = 'EAP', mean=.5))

# custom MAP prior: standard truncated normal between 5 and -2
library(msm)
# need the :: scope for parallel to see the function (not require if no mirtCluster() defined)
fun <- function(Theta, ...) msm::dtnorm(Theta, mean = 0, sd = 1, lower = -2, upper = 5)
head(fscores(mod, custom_den = fun, method = 'MAP', full.scores = FALSE))

#####

```

```

# scoring via response.pattern input (with latent regression structure)
# simulate data
set.seed(1234)
N <- 1000

# covariates
X1 <- rnorm(N); X2 <- rnorm(N)
covdata <- data.frame(X1, X2)
Theta <- matrix(0.5 * X1 + -1 * X2 + rnorm(N, sd = 0.5))

# items and response data
a <- matrix(1, 20); d <- matrix(rnorm(20))
dat <- simdata(a, d, 1000, itemtype = '2PL', Theta=Theta)

# conditional model using X1 and X2 as predictors of Theta
mod <- mirt(dat, 1, 'Rasch', covdata=covdata, formula = ~ X1 + X2)
coef(mod, simplify=TRUE)

# all EAP estimates that include latent regression information
fs <- fscores(mod, full.scores.SE=TRUE)
head(fs)

# score only two response patterns
rp <- dat[1:2, ]
cd <- covdata[1:2, ]

fscores(mod, response.pattern=rp, covdata=cd)
fscores(mod, response.pattern=rp[2,], covdata=cd[2,]) # just one pattern

## End(Not run)

```

gen.difficulty

Generalized item difficulty summaries

Description

Function provides the four generalized item difficulty representations for polytomous response models described by Ali, Chang, and Anderson (2015). These estimates are used to gauge how difficult a polytomous item may be.

Usage

```
gen.difficulty(mod, type = "IRF", interval = c(-30, 30), ...)
```

Arguments

mod a single factor model estimated by [mirt](#)

type	type of generalized difficulty parameter to report. Can be 'IRF' to use the item response function (default), 'mean' to find the average of the difficulty estimates, 'median' the median of the difficulty estimates, and 'trimmed' to find the trimmed mean after removing the first and last difficulty estimates
interval	interval range to search for 'IRF' type
...	additional arguments to pass to uniroot

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Ali, U. S., Chang, H.-H., & Anderson, C. J. (2015). *Location indices for ordinal polytomous items based on item response theory* (Research Report No. RR-15-20). Princeton, NJ: Educational Testing Service. <http://dx.doi.org/10.1002/ets2.12065>

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
## Not run:

mod <- mirt(Science, 1)
coef(mod, simplify=TRUE, IRTpars = TRUE)$items

gen.difficulty(mod)
gen.difficulty(mod, type = 'mean')

# also works for dichotomous items (though this is unnecessary)
dat <- expand.table(LSAT7)
mod <- mirt(dat, 1)
coef(mod, simplify=TRUE, IRTpars = TRUE)$items

gen.difficulty(mod)
gen.difficulty(mod, type = 'mean')

## End(Not run)
```

Description

Given an estimated model from any of mirt's model fitting functions and an estimate of the latent trait, impute plausible missing data values. Returns the original data in a data.frame without any NA values. If a list of Theta values is supplied then a list of complete datasets is returned instead.

Usage

```
imputeMissing(x, Theta, warn = TRUE, ...)
```

Arguments

x	an estimated model x from the mirt package
Theta	a matrix containing the estimates of the latent trait scores (e.g., via fscores)
warn	logical; print warning messages?
...	additional arguments to pass

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
## Not run:
dat <- expand.table(LSAT7)
(original <- mirt(dat, 1))
NAperson <- sample(1:nrow(dat), 20, replace = TRUE)
NAitem <- sample(1:ncol(dat), 20, replace = TRUE)
for(i in 1:20)
  dat[NAperson[i], NAitem[i]] <- NA
(mod <- mirt(dat, 1))
scores <- fscores(mod, method = 'MAP')

# re-estimate imputed dataset (good to do this multiple times and average over)
fulldata <- imputeMissing(mod, scores)
(fullmod <- mirt(fulldata, 1))

# with multipleGroup
set.seed(1)
group <- sample(c('group1', 'group2'), 1000, TRUE)
mod2 <- multipleGroup(dat, 1, group, TOL=1e-2)
fs <- fscores(mod2)
fulldata2 <- imputeMissing(mod2, fs)

## End(Not run)
```

itemfit

*Item fit statistics***Description**

Computes item-fit statistics for a variety of unidimensional and multidimensional models. Poorly fitting items should be inspected with the empirical plots/tables for unidimensional models, otherwise `itemGAM` can be used to diagnose where the functional form of the IRT model was misspecified, or models can be refit using more flexible semi-parametric response models (e.g., `itemtype = 'spline'`). If the latent trait density was approximated (e.g., Davidian curves, Empirical histograms, etc) then passing `use_dentype_estimate = TRUE` will use the internally saved quadrature and density components (where applicable). Currently, only S-X2 statistic supported for mixture IRT models. Finally, where applicable the root mean-square error of approximation (RMSEA) is reported to help gauge the magnitude of item misfit.

Usage

```

itemfit(
  x,
  fit_stats = "S_X2",
  which.items = 1:extract.mirt(x, "nitems"),
  na.rm = FALSE,
  p.adjust = "none",
  group.bins = 10,
  group.size = NA,
  group.fun = mean,
  mincell = 1,
  mincell.X2 = 2,
  return.tables = FALSE,
  pv_draws = 30,
  boot = 1000,
  boot_dfapprox = 200,
  S_X2.plot = NULL,
  S_X2.plot_raw.score = TRUE,
  ETrange = c(-2, 2),
  ETpoints = 11,
  empirical.plot = NULL,
  empirical.CI = 0.95,
  empirical.poly.collapse = FALSE,
  method = "EAP",
  Theta = NULL,
  par.strip.text = list(cex = 0.7),
  par.settings = list(strip.background = list(col = "#9ECAE1"), strip.border = list(col =
    "black")),
  auto.key = list(space = "right", points = FALSE, lines = TRUE),
  ...
)
```

Arguments

<code>x</code>	a computed model object of class <code>SingleGroupClass</code> , <code>MultipleGroupClass</code> , or <code>DiscreteClass</code>
<code>fit_stats</code>	a character vector indicating which fit statistics should be computed. Supported inputs are: • <code>'S_X2'</code> : Orlando and Thissen (2000, 2003) and Kang and Chen's (2007) signed chi-squared test (default) • <code>'Zh'</code> : Drasgow, Levine, & Williams (1985) Zh • <code>'X2'</code> : Bock's (1972) chi-squared method. The default inputs compute Yen's (1981) Q1 variant of the X2 statistic (i.e., uses a fixed <code>group.bins = 10</code>). However, Bock's group-size variable median-based method can be computed by passing <code>group.fun = median</code> and modifying the <code>group.size</code> input to the desired number of bins • <code>'G2'</code> : McKinley & Mills (1985) G2 statistic (similar method to Q1, but with the likelihood-ratio test). • <code>'PV_Q1'</code> : Chalmers and Ng's (2017) plausible-value variant of the Q1 statistic. • <code>'PV_Q1*'</code> : Chalmers and Ng's (2017) plausible-value variant of the Q1 statistic that uses parametric bootstrapping to obtain a suitable empirical distribution. • <code>'X2*'</code> : Stone's (2000) fit statistics that require parametric bootstrapping • <code>'X2*_df'</code> : Stone's (2000) fit statistics that require parametric bootstrapping to obtain scaled versions of the X2* and degrees of freedom • <code>'infit'</code> : Compute the infit and outfit statistics Note that <code>'S_X2'</code> and <code>'Zh'</code> cannot be computed when there are missing response data (i.e., will require multiple-imputation/row-removal techniques).
<code>which.items</code>	an integer vector indicating which items to test for fit. Default tests all possible items
<code>na.rm</code>	logical; remove rows with any missing values? This is required for methods such as S-X2 because they require the "EAPsum" method from fscores
<code>p.adjust</code>	method to use for adjusting all p-values for each respective item fit statistic (see p.adjust for available options). Default is <code>'none'</code>
<code>group.bins</code>	the number of bins to use for X2 and G2. For example, setting <code>group.bins = 10</code> will compute Yen's (1981) Q1 statistic when <code>'X2'</code> is requested
<code>group.size</code>	approximate size of each group to be used in calculating the χ^2 statistic. The default NA disables this command and instead uses the <code>group.bins</code> input to try and construct equally sized bins
<code>group.fun</code>	function used when <code>'X2'</code> or <code>'G2'</code> are computed. Determines the central tendency measure within each partitioned group. E.g., setting <code>group.fun = median</code> will obtain the median of each respective ability estimate in each subgroup (this is what was used by Bock, 1972)
<code>mincell</code>	the minimum expected cell size to be used in the S-X2 computations. Tables will be collapsed across items first if polytomous, and then across scores if necessary

<code>mincell.X2</code>	the minimum expected cell size to be used in the X2 computations. Tables will be collapsed if polytomous, however if this condition can not be met then the group block will be omitted in the computations
<code>return.tables</code>	logical; return tables when investigating 'X2', 'S_X2', and 'X2*'?
<code>pv_draws</code>	number of plausible-value draws to obtain for PV_Q1 and PV_Q1*
<code>boot</code>	number of parametric bootstrap samples to create for PV_Q1* and X2*
<code>boot_dfapprox</code>	number of parametric bootstrap samples to create for the X2*_df statistic to approximate the scaling factor for X2* as well as the scaled degrees of freedom estimates
<code>S_X2.plot</code>	argument input is the same as <code>empirical.plot</code> , however the resulting image is constructed according to the S-X2 statistic's conditional sum-score information
<code>S_X2.plot_raw.score</code>	logical; use the raw-score information in the plot in stead of the latent trait scale score? Default is FALSE
<code>ETrange</code>	range of integration nodes for Stone's X2* statistic
<code>ETpoints</code>	number of integration nodes to use for Stone's X2* statistic
<code>empirical.plot</code>	a single numeric value or character of the item name indicating which item to plot (via <code>itemplot</code>) and overlay with the empirical θ groupings (see <code>empirical.CI</code>). Useful for plotting the expected bins based on the 'X2' or 'G2' method
<code>empirical.CI</code>	a numeric value indicating the width of the empirical confidence interval ranging between 0 and 1 (default of 0 plots not interval). For example, a 95 interval would be plotted when <code>empirical.CI</code> = .95. Only applicable to dichotomous items
<code>empirical.poly.collapse</code>	logical; collapse polytomous item categories to for expected scoring functions for empirical plots? Default is FALSE
<code>method</code>	type of factor score estimation method. See fscores for more detail
<code>Theta</code>	a matrix of factor scores for each person used for statistics that require empirical estimates. If supplied, arguments typically passed to <code>fscores()</code> will be ignored and these values will be used instead. Also required when estimating statistics with missing data via imputation
<code>par.strip.text</code>	plotting argument passed to lattice
<code>par.settings</code>	plotting argument passed to lattice
<code>auto.key</code>	plotting argument passed to lattice
<code>...</code>	additional arguments to be passed to <code>fscores()</code> and lattice

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Bock, R. D. (1972). Estimating item parameters and latent ability when responses are scored in two or more nominal categories. *Psychometrika*, 37, 29-51.
- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Chalmers, R. P. & Ng, V. (2017). Plausible-Value Imputation Statistics for Detecting Item Misfit. *Applied Psychological Measurement*, 41, 372-387. doi:10.1177/0146621617692079
- Drasgow, F., Levine, M. V., & Williams, E. A. (1985). Appropriateness measurement with polychotomous item response models and standardized indices. *British Journal of Mathematical and Statistical Psychology*, 38, 67-86.
- Kang, T. & Chen, Troy, T. (2007). An investigation of the performance of the generalized S-X2 item-fit index for polytomous IRT models. ACT
- McKinley, R., & Mills, C. (1985). A comparison of several goodness-of-fit statistics. *Applied Psychological Measurement*, 9, 49-57.
- Orlando, M. & Thissen, D. (2000). Likelihood-based item fit indices for dichotomous item response theory models. *Applied Psychological Measurement*, 24, 50-64.
- Reise, S. P. (1990). A comparison of item- and person-fit methods of assessing model-data fit in IRT. *Applied Psychological Measurement*, 14, 127-137.
- Stone, C. A. (2000). Monte Carlo Based Null Distribution for an Alternative Goodness-of-Fit Test Statistics in IRT Models. *Journal of Educational Measurement*, 37, 58-75.
- Wright B. D. & Masters, G. N. (1982). *Rating scale analysis*. MESA Press.
- Yen, W. M. (1981). Using simulation results to choose a latent trait model. *Applied Psychological Measurement*, 5, 245-262.

See Also

[personfit](#), [itemGAM](#)

Examples

```
## Not run:

P <- function(Theta){exp(Theta^2 * 1.2 - 1) / (1 + exp(Theta^2 * 1.2 - 1))}

#make some data
set.seed(1234)
a <- matrix(rlnorm(20, meanlog=0, sdlog = .1),ncol=1)
d <- matrix(rnorm(20),ncol=1)
Theta <- matrix(rnorm(2000))
items <- rep('2PL', 20)
ps <- P(Theta)
baditem <- numeric(2000)
for(i in 1:2000)
  baditem[i] <- sample(c(0,1), 1, prob = c(1-ps[i], ps[i]))
data <- cbind(simdata(a,d, 2000, items, Theta=Theta), baditem=baditem)

x <- mirt(data, 1)
```

```

raschfit <- mirt(data, 1, itemtype='Rasch')
fit <- itemfit(x)
fit

# p-value adjustment
itemfit(x, p.adjust='fdr')

# two different fit stats (with/without p-value adjustment)
itemfit(x, c('S_X2' , 'X2'), p.adjust='fdr')
itemfit(x, c('S_X2' , 'X2'))

# Conditional sum-score plot from S-X2 information
itemfit(x, S_X2.plot = 1) # good fit
itemfit(x, S_X2.plot = 2) # good fit
itemfit(x, S_X2.plot = 21) # bad fit

itemfit(x, 'X2') # just X2
itemfit(x, 'X2', method = 'ML') # X2 with maximum-likelihood estimates for traits
itemfit(x, group.bins=15, empirical.plot = 1, method = 'ML') #empirical item plot with 15 points
itemfit(x, group.bins=15, empirical.plot = 21, method = 'ML')

# PV and X2* statistics (parametric bootstrap stats not run to save time)
itemfit(x, 'PV_Q1')

if(interactive()) mirtCluster() # improve speed of bootstrap samples by running in parallel
# itemfit(x, 'PV_Q1*')
# itemfit(x, 'X2*') # Stone's 1993 statistic
# itemfit(x, 'X2*_df') # Stone's 2000 scaled statistic with df estimate

# empirical tables for X2 statistic
tabs <- itemfit(x, 'X2', return.tables=TRUE, which.items = 1)
tabs

#infit/outfit statistics. method='ML' agrees better with eRm package
itemfit(raschfit, 'infit', method = 'ML') #infit and outfit stats

#same as above, but inputting ML estimates instead (saves time for re-use)
Theta <- fscores(raschfit, method = 'ML')
itemfit(raschfit, 'infit', Theta=Theta)
itemfit(raschfit, empirical.plot=1, Theta=Theta)
itemfit(raschfit, 'X2', return.tables=TRUE, Theta=Theta, which.items=1)

# fit a new more flexible model for the mis-fitting item
itemtype <- c(rep('2PL', 20), 'spline')
x2 <- mirt(data, 1, itemtype=itemtype)
itemfit(x2)
itemplot(x2, 21)
anova(x, x2)

#-----

#similar example to Kang and Chen 2007
a <- matrix(c(.8,.4,.7, .8, .4, .7, 1, 1, 1, 1))

```

```

d <- matrix(rep(c(2.0,0.0,-1,-1.5),10), ncol=4, byrow=TRUE)
dat <- simdata(a,d,2000, itemtype = rep('graded', 10))
head(dat)

mod <- mirt(dat, 1)
itemfit(mod)
itemfit(mod, 'X2') # less useful given inflated Type I error rates
itemfit(mod, empirical.plot = 1)
itemfit(mod, empirical.plot = 1, empirical.poly.collapse=TRUE)

# collapsed tables (see mincell.X2) for X2 and G2
itemfit(mod, 'X2', return.tables = TRUE, which.items = 1)

mod2 <- mirt(dat, 1, 'Rasch')
itemfit(mod2, 'infit', method = 'ML')

# massive list of tables for S-X2
tables <- itemfit(mod, return.tables = TRUE)

#observed and expected total score patterns for item 1 (post collapsing)
tables$O[[1]]
tables$E[[1]]

# can also select specific items
# itemfit(mod, return.tables = TRUE, which.items=1)

# fit stats with missing data (run in parallel using all cores)
dat[sample(1:prod(dim(dat)), 100)] <- NA
raschfit <- mirt(dat, 1, itemtype='Rasch')

# use only valid data by removing rows with missing terms
itemfit(raschfit, c('S_X2', 'infit'), na.rm = TRUE)

# note that X2, G2, PV-Q1, and X2* do not require complete datasets
thetas <- fscores(raschfit, method = 'ML') # save for faster computations
itemfit(raschfit, c('X2', 'G2'), Theta=thetas)
itemfit(raschfit, empirical.plot=1, Theta=thetas)
itemfit(raschfit, 'X2', return.tables=TRUE, which.items=1, Theta=thetas)

## End(Not run)

```

itemGAM

Parametric smoothed regression lines for item response probability functions

Description

This function uses a generalized additive model (GAM) to estimate response curves for items that do not seem to fit well in a given model. Using a stable axillary model, traceline functions for poorly

fitting dichotomous or polytomous items can be inspected using point estimates (or plausible values) of the latent trait. Plots of the tracelines and their associated standard errors are available to help interpret the misfit. This function may also be useful when adding new items to an existing, well established set of items, especially when the parametric form of the items under investigation are unknown.

Usage

```
itemGAM(
  item,
  Theta,
  formula = resp ~ s(Theta, k = 10),
  CI = 0.95,
  theta_lim = c(-3, 3),
  return.models = FALSE,
  ...
)

## S3 method for class 'itemGAM'
plot(
  x,
  y = NULL,
  par.strip.text = list(cex = 0.7),
  par.settings = list(strip.background = list(col = "#9ECAE1"), strip.border = list(col =
    "black")),
  auto.key = list(space = "right", points = FALSE, lines = TRUE),
  ...
)
```

Arguments

<code>item</code>	a single poorly fitting item to be investigated. Can be a vector or matrix
<code>Theta</code>	a list or matrix of latent trait estimates typically returned from fscores
<code>formula</code>	an R formula to be passed to the <code>gam</code> function. Default fits a spline model with 10 nodes. For multidimensional models, the traits are assigned the names 'Theta1', 'Theta2', ..., 'ThetaN'
<code>CI</code>	a number ranging from 0 to 1 indicating the confidence interval range. Default provides the 95 percent interval
<code>theta_lim</code>	range of latent trait scores to be evaluated
<code>return.models</code>	logical; return a list of GAM models for each category? Useful when the GAMs should be inspected directly, but also when fitting multidimensional models (this is set to TRUE automatically for multidimensional models)
<code>...</code>	additional arguments to be passed to <code>gam</code> or <code>lattice</code>
<code>x</code>	an object of class 'itemGAM'
<code>y</code>	a NULL value ignored by the plotting function
<code>par.strip.text</code>	plotting argument passed to lattice

par.settings plotting argument passed to [lattice](#)
 auto.key plotting argument passed to [lattice](#)

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[itemfit](#)

Examples

```
## Not run:
set.seed(10)
N <- 1000
J <- 30

a <- matrix(1, J)
d <- matrix(rnorm(J))
Theta <- matrix(rnorm(N, 0, 1.5))
dat <- simdata(a, d, N, itemtype = '2PL', Theta=Theta)

# make a bad item
ps <- exp(Theta^2 + Theta) / (1 + exp(Theta^2 + Theta))
item1 <- sapply(ps, function(x) sample(c(0,1), size = 1, prob = c(1-x, x)))

ps2 <- exp(2 * Theta^2 + Theta + .5 * Theta^3) / (1 + exp(2 * Theta^2 + Theta + .5 * Theta^3))
item2 <- sapply(ps2, function(x) sample(c(0,1), size = 1, prob = c(1-x, x)))

# how the actual item looks in the population
plot(Theta, ps, ylim = c(0,1))
plot(Theta, ps2, ylim = c(0,1))

baditems <- cbind(item1, item2)
newdat <- cbind(dat, baditems)

badmod <- mirt(newdat, 1)
itemfit(badmod) #clearly a bad fit for the last two items
mod <- mirt(dat, 1) #fit a model that does not contain the bad items
itemfit(mod)

#### Pure non-parametric way of investigating the items
library(KernSmoothIRT)
ks <- ksIRT(newdat, rep(1, ncol(newdat)), 1)
plot(ks, item=c(1,31,32))
par(ask=FALSE)
```

```

# Using point estimates from the model
Theta <- fscores(mod)
IG0 <- itemGAM(dat[,1], Theta) #good item
IG1 <- itemGAM(baditems[,1], Theta)
IG2 <- itemGAM(baditems[,2], Theta)
plot(IG0)
plot(IG1)
plot(IG2)

# same as above, but with plausible values to obtain the standard errors
set.seed(4321)
ThetaPV <- fscores(mod, plausible.draws=10)
IG0 <- itemGAM(dat[,1], ThetaPV) #good item
IG1 <- itemGAM(baditems[,1], ThetaPV)
IG2 <- itemGAM(baditems[,2], ThetaPV)
plot(IG0)
plot(IG1)
plot(IG2)

## for polytomous test items
SAT12[SAT12 == 8] <- NA
dat <- key2binary(SAT12,
                  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
dat <- dat[,~32]
mod <- mirt(dat, 1)

# Kernal smoothing is very sensitive to which category is selected as 'correct'
# 5th category as correct
ks <- ksIRT(cbind(dat, SAT12[,32]), c(rep(1, 31), 5), 1)
plot(ks, items = c(1,2,32))

# 3rd category as correct
ks <- ksIRT(cbind(dat, SAT12[,32]), c(rep(1, 31), 3), 1)
plot(ks, items = c(1,2,32))

# splines approach
Theta <- fscores(mod)
IG <- itemGAM(SAT12[,32], Theta)
plot(IG)

set.seed(1423)
ThetaPV <- fscores(mod, plausible.draws=10)
IG2 <- itemGAM(SAT12[,32], ThetaPV)
plot(IG2)

# assuming a simple increasing parametric form (like in a standard IRT model)
IG3 <- itemGAM(SAT12[,32], Theta, formula = resp ~ Theta)
plot(IG3)
IG3 <- itemGAM(SAT12[,32], ThetaPV, formula = resp ~ Theta)
plot(IG3)

### multidimensional example by returning the GAM objects

```

```

mod2 <- mirt(dat, 2)
Theta <- fscores(mod2)
IG4 <- itemGAM(SAT12[,32], Theta, formula = resp ~ s(Theta1, k=10) + s(Theta2, k=10),
  return.models=TRUE)
names(IG4)
plot(IG4[[1L]], main = 'Category 1')
plot(IG4[[2L]], main = 'Category 2')
plot(IG4[[3L]], main = 'Category 3')

## End(Not run)

```

iteminfo

Function to calculate item information

Description

Given an internal mirt item object extracted by using [extract.item](#), compute the item information.

Usage

```
iteminfo(x, Theta, degrees = NULL, total.info = TRUE, multidim_matrix = FALSE)
```

Arguments

x	an extracted internal mirt object containing item information (see extract.item)
Theta	a vector (unidimensional) or matrix (multidimensional) of latent trait values
degrees	a vector of angles in degrees that are between 0 and 90. Only applicable when the input object is multidimensional
total.info	logical; return the total information curve for the item? If FALSE, information curves for each category are returned as a matrix
multidim_matrix	logical; compute the information matrix for each row in Theta? If Theta contains more than 1 row then a list of matrices will be returned, otherwise if Theta has exactly one row then a matrix will be returned

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[extract.item](#)

Examples

```

mod <- mirt(Science, 1)
extr.2 <- extract.item(mod, 2)
Theta <- matrix(seq(-4,4, by = .1))
info.2 <- iteminfo(extr.2, Theta)

#do something with the info?
plot(Theta, info.2, type = 'l', main = 'Item information')

## Not run:

#category information curves
cat.info <- iteminfo(extr.2, Theta, total.info = FALSE)
plot(Theta, cat.info[,1], type = 'l', ylim = c(0, max(cat.info)),
     ylab = 'info', main = 'Category information')
for(i in 2:ncol(cat.info))
  lines(Theta, cat.info[,i], col = i)

## Customized test information plot
T1 <- T2 <- 0
dat <- expand.table(LSAT7)
mod1 <- mirt(dat, 1)
mod2 <- mirt(dat, 1, 'Rasch')
for(i in 1:5){
  T1 <- T1 + iteminfo(extract.item(mod1, i), Theta)
  T2 <- T2 + iteminfo(extract.item(mod2, i), Theta)
}
plot(Theta, T2/T1, type = 'l', ylab = 'Relative Test Information', las = 1)
lines(Theta, T1/T1, col = 'red')

# multidimensional
mod <- mirt(dat, 2, TOL=1e-2)
ii <- extract.item(mod, 1)
Theta <- as.matrix(expand.grid(-4:4, -4:4))

iteminfo(ii, Theta, degrees=c(45,45)) # equal angle
iteminfo(ii, Theta, degrees=c(90,0)) # first dimension only

# information matrices
iteminfo(ii, Theta, multidim_matrix = TRUE)
iteminfo(ii, Theta[1, , drop=FALSE], multidim_matrix = TRUE)

## End(Not run)

```

Description

itemplot displays various item based IRT plots, with special options for plotting items that contain several 0 slope parameters. Supports up to three dimensional models.

Usage

```
itemplot(
  object,
  item,
  type = "trace",
  degrees = 45,
  CE = FALSE,
  CEalpha = 0.05,
  CEdraws = 1000,
  drop.zeros = FALSE,
  theta_lim = c(-6, 6),
  shiny = FALSE,
  rot = list(xaxis = -70, yaxis = 30, zaxis = 10),
  par.strip.text = list(cex = 0.7),
  npts = 200,
  par.settings = list(strip.background = list(col = "#9ECAE1"), strip.border = list(col =
    "black")),
  auto.key = list(space = "right", points = FALSE, lines = TRUE),
  ...
)
```

Arguments

object	a computed model object of class SingleGroupClass or MultipleGroupClass. Input may also be a list for comparing similar item types (e.g., 1PL vs 2PL)
item	a single numeric value, or the item name, indicating which item to plot
type	plot type to use, information ('info'), standard errors ('SE'), item trace lines ('trace'), cumulative probability plots to indicate thresholds ('threshold'), information and standard errors ('infoSE') or information and trace lines ('infotrace'), category and total information ('infocat'), relative efficiency lines ('RE'), expected score 'score', or information and trace line contours ('infocontour' and 'tracecontour'; not supported for MultipleGroupClass objects)
degrees	the degrees argument to be used if there are two or three factors. See iteminfo for more detail. A new vector will be required for three dimensional models to override the default
CE	logical; plot confidence envelope?
CEalpha	area remaining in the tail for confidence envelope. Default gives 95% confidence region
CEdraws	draws number of draws to use for confidence envelope
drop.zeros	logical; drop slope values that are numerically close to zero to reduce dimensionality? Useful in objects returned from bfactor or other confirmatory models that contain several zero slopes

theta_lim	lower and upper limits of the latent trait (theta) to be evaluated, and is used in conjunction with npts. Default uses <code>c(-6, 6)</code>
shiny	logical; run interactive display for item plots using the shiny interface. This primarily is an instructive tool for demonstrating how item response curves behave when adjusting their parameters
rot	a list of rotation coordinates to be used for 3 dimensional plots
par.strip.text	plotting argument passed to <code>lattice</code>
npts	number of quadrature points to be used for plotting features. Larger values make plots look smoother
par.settings	plotting argument passed to <code>lattice</code>
auto.key	plotting argument passed to <code>lattice</code>
...	additional arguments to be passed to <code>lattice</code> and <code>coef()</code>

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
## Not run:

data(LSAT7)
fulldata <- expand.table(LSAT7)
mod1 <- mirt(fulldata,1,SE=TRUE)
mod2 <- mirt(fulldata,1, itemtype = 'Rasch')
mod3 <- mirt(fulldata,2)

itemplot(mod1, 2)
itemplot(mod1, 2, CE = TRUE)
itemplot(mod1, 2, type = 'info')
itemplot(mod1, 2, type = 'info', CE = TRUE)

mods <- list(twoPL = mod1, onePL = mod2)
itemplot(mods, 1, type = 'RE')

# multidimensional
itemplot(mod3, 4, type = 'info')
itemplot(mod3, 4, type = 'info',
  col.regions = colorRampPalette(c("white", "red"))(100))
itemplot(mod3, 4, type = 'infocontour')
itemplot(mod3, 4, type = 'tracecontour')

# polytomous items
pmod <- mirt(Science, 1, SE=TRUE)
```

```

itemplot(pmod, 3)
itemplot(pmod, 3, type = 'threshold')
itemplot(pmod, 3, CE = TRUE)
itemplot(pmod, 3, type = 'score')
itemplot(pmod, 3, type = 'score', CE = TRUE)
itemplot(pmod, 3, type = 'infotrace')
itemplot(pmod, 3, type = 'infocat')

# use the directlabels package to put labels on tracelines
library(directlabels)
plt <- itemplot(pmod, 3)
direct.label(plt, 'top.points')

# change colour theme of plots
bwtheme <- standard.theme("pdf", color=FALSE)
plot(pmod, type='trace', par.settings=bwtheme)
itemplot(pmod, 1, type = 'trace', par.settings=bwtheme)

# additional modifications can be made via update().
# See ?update.trellis for further documentation
(plt <- itemplot(pmod, 1))
update(plt, ylab = expression(Prob(theta))) # ylab changed

# infoSE plot
itemplot(pmod, 1, type = 'infoSE')

# uncomment to run interactive shiny applet
# itemplot(shiny = TRUE)

## End(Not run)

```

itemstats

Generic item summary statistics

Description

Function to compute generic item summary statistics that do not require prior fitting of IRT models. Contains information about coefficient alpha (and alpha if an item is deleted), mean/SD and frequency of total scores, reduced item-total correlations, average/sd of the correlation between items, response frequencies, and conditional mean/sd information given the unweighted sum scores.

Usage

```

itemstats(
  data,
  group = NULL,
  use_ts = TRUE,
  proportions = TRUE,

```

```
    ts.tables = FALSE
  )
```

Arguments

data	An object of class <code>data.frame</code> or <code>matrix</code> with the response patterns
group	optional grouping variable to condition on when computing summary information
use_ts	logical; include information that is conditional on a meaningful total score?
proportions	logical; include response proportion information for each item?
ts.tables	logical; include mean/sd summary information pertaining to the unweighted total score?

Value

Returns a list containing the summary statistics

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[empirical_plot](#)

Examples

```
# dichotomous data example
LSAT7full <- expand.table(LSAT7)
head(LSAT7full)
itemstats(LSAT7full)

# behaviour with missing data
LSAT7full[1:5,1] <- NA
itemstats(LSAT7full)

# data with no meaningful total score
head(SAT12)
itemstats(SAT12, use_ts=FALSE)

# extra total scores tables
dat <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,
          5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
itemstats(dat, ts.tables=TRUE)
```

```

# grouping information
group <- gl(2, 300, labels=c('G1', 'G2'))
itemstats(dat, group=group)

#####
# polytomous data example
itemstats(Science)

# polytomous data with missing
newScience <- Science
newScience[1:5,1] <- NA
itemstats(newScience)

# unequal categories
newScience[,1] <- ifelse(Science[,1] == 1, NA, Science[,1])
itemstats(newScience)

merged <- data.frame(LSAT7full[1:392,], Science)
itemstats(merged)

```

key2binary

Score a test by converting response patterns to binary data

Description

The key2binary function will convert response pattern data to a dichotomous format, given a response key.

Usage

```
key2binary(fulldata, key, score_missing = FALSE)
```

Arguments

fulldata	an object of class <code>data.frame</code> , <code>matrix</code> , or <code>table</code> with the response patterns
key	a vector or matrix consisting of the 'correct' response to the items. Each value/row corresponds to each column in <code>fulldata</code> . If the input is a matrix, multiple scoring keys can be supplied for each item. NA values are used to indicate no scoring key (or in the case of a matrix input, no additional scoring keys)
score_missing	logical; should missing data elements be returned as incorrect (i.e., 0)? If FALSE, all missing data terms will be kept as missing

Value

Returns a numeric matrix with all the response patterns in dichotomous format

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
data(SAT12)
head(SAT12)
key <- c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5)

dicho.SAT12 <- key2binary(SAT12, key)
head(dicho.SAT12)

# multiple scoring keys
key2 <- cbind(c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5),
              c(2,3,NA,1,rep(NA, 28)))
dicho.SAT12 <- key2binary(SAT12, key2)

# keys from raw character responses
resp <- as.data.frame(matrix(c(
  "B","B","D","D","E",
  "B","A","D","D","E",
  "B","A","D","C","E",
  "D","D","D","C","E",
  "B","C","A","D","A"), ncol=5, byrow=TRUE))

key <- c("B", "D", "D", "C", "E")

d01 <- key2binary(resp, key)
head(d01)

# score/don't score missing values
resp[1,1] <- NA
d01NA <- key2binary(resp, key) # without scoring
d01NA

d01 <- key2binary(resp, key, score_missing = TRUE) # with scoring
d01
```

Description

Lagrange (i.e., score) test to test whether parameters should be freed from a more constrained baseline model.

Usage

```
lagrange(mod, parnum, SE.type = "Oakes", type = "Richardson", ...)
```

Arguments

<code>mod</code>	an estimated model
<code>parnum</code>	a vector, or list of vectors, containing one or more parameter locations/sets of locations to be tested. See objects returned from mod2values for the locations
<code>SE.type</code>	type of information matrix estimator to use. See mirt for further details
<code>type</code>	type of numerical algorithm passed to numerical_deriv to obtain the gradient terms
<code>...</code>	additional arguments to pass to mirt

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](#)

See Also

[wald](#)

Examples

```
## Not run:
dat <- expand.table(LSAT7)
mod <- mirt(dat, 1, 'Rasch')
(values <- mod2values(mod))

# test all fixed slopes individually
parnum <- values$parnum[values$name == 'a1']
lagrange(mod, parnum)

# compare to LR test for first two slopes
mod2 <- mirt(dat, 'F = 1-5
              FREE = (1, a1)', 'Rasch')
coef(mod2, simplify=TRUE)$items
anova(mod, mod2)

mod2 <- mirt(dat, 'F = 1-5
```

```

                                FREE = (2, a1)', 'Rasch')
coef(mod2, simplify=TRUE)$items
anova(mod, mod2)

mod2 <- mirt(dat, 'F = 1-5
                                FREE = (3, a1)', 'Rasch')
coef(mod2, simplify=TRUE)$items
anova(mod, mod2)

# test slopes first two slopes and last three slopes jointly
lagrange(mod, list(parnum[1:2], parnum[3:5]))

# test all 5 slopes and first + last jointly
lagrange(mod, list(parnum[1:5], parnum[c(1, 5)]))

## End(Not run)

```

likert2int	<i>Convert ordered Likert-scale responses (character or factors) to integers</i>
------------	--

Description

Given a matrix or data.frame object consisting of Likert responses return an object of the same dimensions with integer values.

Usage

```
likert2int(x, levels = NULL)
```

Arguments

x	a matrix of character values or data.frame of character/factor vectors
levels	a named character vector indicating which integer values should be assigned to which elements. If omitted, the order of the elements will be determined after converting each column in x to a factor variable

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[key2binary](#), [poly2dich](#)

Examples

```
## Not run:

# simulate data

dat1 <- matrix(sample(c('Disagree', 'Strongly Disagree', 'Agree',
                        'Neutral', 'Strongly Agree'), 1000*5, replace=TRUE),
               nrow=1000, ncol=5)
dat1[2,2] <- dat1[3,3] <- dat1[1,3] <- NA # NAs added for flavour
dat2 <- matrix(sample(c('D', 'SD', 'A', 'N', 'SA'), 1000*5, replace=TRUE),
               nrow=1000, ncol=5)
dat <- cbind(dat1, dat2)

# separately
intdat1 <- likert2int(dat1)
head(dat1)
head(intdat1)

# more useful with explicit levels
lv11 <- c('Strongly Disagree'=1, 'Disagree'=2, 'Neutral'=3, 'Agree'=4,
          'Strongly Agree'=5)
intdat1 <- likert2int(dat1, levels = lv11)
head(dat1)
head(intdat1)

# second data
lv12 <- c('SD'=1, 'D'=2, 'N'=3, 'A'=4, 'SA'=5)
intdat2 <- likert2int(dat2, levels = lv12)
head(dat2)
head(intdat2)

# full dataset (using both mapping schemes)
intdat <- likert2int(dat, levels = c(lv11, lv12))
head(dat)
head(intdat)

#####
# data.frame as input with ordered factors

dat1 <- data.frame(dat1)
dat2 <- data.frame(dat2)
dat.old <- cbind(dat1, dat2)
colnames(dat.old) <- paste0('Item_', 1:10)
str(dat.old) # factors are leveled alphabetically by default

# create explicit ordering in factor variables
for(i in 1:ncol(dat1))
  levels(dat1[[i]]) <- c('Strongly Disagree', 'Disagree', 'Neutral', 'Agree',
                        'Strongly Agree')

for(i in 1:ncol(dat2))
```

```

      levels(dat2[[i]]) <- c('SD', 'D', 'N', 'A', 'SA')

  dat <- cbind(dat1, dat2)
  colnames(dat) <- colnames(dat.old)
  str(dat) # note ordering

  intdat <- likert2int(dat)
  head(dat)
  head(intdat)

## End(Not run)

```

logLik-method

Extract log-likelihood

Description

Extract the observed-data log-likelihood.

Usage

```

## S4 method for signature 'SingleGroupClass'
logLik(object)

```

Arguments

object an object of class `SingleGroupClass`, `MultipleGroupClass`, or `MixedClass`

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```

## Not run:
x <- mirt(Science, 1)
logLik(x)

## End(Not run)

```

LSAT6

*Description of LSAT6 data***Description**

Data from Thissen (1982); contains 5 dichotomously scored items obtained from the Law School Admissions Test, section 6.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Thissen, D. (1982). Marginal maximum likelihood estimation for the one-parameter logistic model. *Psychometrika*, 47, 175-186.

Examples

```
## Not run:
dat <- expand.table(LSAT6)
head(dat)
itemstats(dat)

model <- 'F = 1-5
          CONSTRAIN = (1-5, a1)'
(mod <- mirt(dat, model))
M2(mod)
itemfit(mod)
coef(mod, simplify=TRUE)

# equivalently, but with a different parameterization
mod2 <- mirt(dat, 1, itemtype = 'Rasch')
anova(mod, mod2) #equal
M2(mod2)
itemfit(mod2)
coef(mod2, simplify=TRUE)
sqrt(coef(mod2)$GroupPars[2]) #latent SD equal to the slope in mod

## End(Not run)
```

LSAT7*Description of LSAT7 data*

Description

Data from Bock & Lieberman (1970); contains 5 dichotomously scored items obtained from the Law School Admissions Test, section 7.

Data from

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Bock, R. D., & Lieberman, M. (1970). Fitting a response model for n dichotomously scored items. *Psychometrika*, 35(2), 179-197.

Bock, R. D., & Lieberman, M. (1970). Fitting a response model for n dichotomously scored items. *Psychometrika*, 35(2), 179-197.

Examples

```
## Not run:  
dat <- expand.table(LSAT7)  
head(dat)  
itemstats(dat)
```

```
(mod <- mirt(dat, 1))  
coef(mod)
```

```
## End(Not run)
```

```
## Not run:  
dat <- expand.table(LSAT7)  
head(dat)  
itemstats(dat)
```

```
(mod <- mirt(dat, 1))  
coef(mod)
```

```
## End(Not run)
```

M2

*Compute the M2 model fit statistic***Description**

Computes the M2 (Maydeu-Olivares & Joe, 2006) statistic when all data are dichotomous, the collapsed M2* statistic (collapsing over univariate and bivariate response categories; see Cai and Hansen, 2013), and the hybrid C2 statistic which only collapses only the bivariate moments (Cai and Monro, 2014). The C2 variant is mainly useful when polytomous response models do not have sufficient degrees of freedom to compute M2*. This function also computes associated fit indices that are based on fitting the null model. Supports single and multiple-group models. If the latent trait density was approximated (e.g., Davidian curves, Empirical histograms, etc) then passing `use_dentype_estimate = TRUE` will use the internally saved quadrature and density components (where applicable).

Usage

```
M2(
  obj,
  type = "M2*",
  calcNull = TRUE,
  quadpts = NULL,
  theta_lim = c(-6, 6),
  CI = 0.9,
  residmat = FALSE,
  QMC = FALSE,
  suppress = 1,
  ...
)
```

Arguments

<code>obj</code>	an estimated model object from the mirt package
<code>type</code>	type of fit statistic to compute. Options are "M2", "M2*" for the univariate and bivariate collapsed version of the M2 statistic ("M2" currently limited to dichotomous response data only), and "C2" for a hybrid between M2 and M2* where only the bivariate moments are collapsed
<code>calcNull</code>	logical; calculate statistics for the null model as well? Allows for statistics such as the limited information TLI and CFI. Only valid when items all have a suitable null model (e.g., those created via createItem will not)
<code>quadpts</code>	number of quadrature points to use during estimation. If NULL, a suitable value will be chosen based on the rubric found in fcores
<code>theta_lim</code>	lower and upper range to evaluate latent trait integral for each dimension
<code>CI</code>	numeric value from 0 to 1 indicating the range of the confidence interval for RMSEA. Default returns the 90% interval

<code>residmat</code>	logical; return the residual matrix used to compute the SRMSR statistic? Only the lower triangle of the residual correlation matrix will be returned (the upper triangle is filled with NA's)
<code>QMC</code>	logical; use quasi-Monte Carlo integration? Useful for higher dimensional models. If <code>quadpts</code> not specified, 5000 nodes are used by default
<code>suppress</code>	a numeric value indicating which parameter residual dependency combinations to flag as being too high. Absolute values for the standardized residuals greater than this value will be returned, while all values less than this value will be set to NA. Must be used in conjunction with the argument <code>residmat = TRUE</code>
<code>...</code>	additional arguments to pass

Value

Returns a data.frame object with the M2-type statistic, along with the degrees of freedom, p-value, RMSEA (with 90% confidence interval), SRMSR for each group, and optionally the TLI and CFI model fit statistics if `calcNull = TRUE`.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Cai, L. & Hansen, M. (2013). Limited-information goodness-of-fit testing of hierarchical item factor models. *British Journal of Mathematical and Statistical Psychology*, 66, 245-276.
- Cai, L. & Monroe, S. (2014). *A new statistic for evaluating item response theory models for ordinal data*. National Center for Research on Evaluation, Standards, & Student Testing. Technical Report.
- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Maydeu-Olivares, A. & Joe, H. (2006). Limited information goodness-of-fit testing in multidimensional contingency tables. *Psychometrika*, 71, 713-732.

Examples

```
## Not run:
dat <- as.matrix(expand.table(LSAT7))
(mod1 <- mirt(dat, 1))
M2(mod1)
resids <- M2(mod1, residmat=TRUE) #lower triangle of residual correlation matrix
resids
summary(resids[lower.tri(resids)])

# M2 with missing data present
dat[sample(1:prod(dim(dat)), 250)] <- NA
mod2 <- mirt(dat, 1)
M2(mod2)

# C2 statistic (useful when polytomous IRT models have too few df)
pmod <- mirt(Science, 1)
```

```
# This fails with too few df:
# M2(pmod)
# This, however, works:
M2(pmod, type = 'C2')

## End(Not run)
```

marginal_rxx

Function to calculate the marginal reliability

Description

Given an estimated model and a prior density function, compute the marginal reliability (Thissen and Wainer, 2001). This is only available for unidimensional tests.

Usage

```
marginal_rxx(mod, density = dnorm, ...)
```

Arguments

mod	an object of class 'SingleGroupClass'
density	a density function to use for integration. Default assumes the latent traits are from a normal (Gaussian) distribution
...	additional arguments passed to the density function

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Thissen, D. and Wainer, H. (2001). Test Scoring. Lawrence Erlbaum Associates.

See Also

[empirical_rxx](#), [extract.group](#), [testinfo](#)

Examples

```

dat <- expand.table(deAyala)
mod <- mirt(dat)

# marginal estimate treating item parameters as known
marginal_rxx(mod)

# compare to alpha
itemstats(dat)$overall$alpha

## Not run:

# empirical estimate (assuming the same prior)
fscores(mod, returnER = TRUE)

# empirical rxx the alternative way, given theta scores and SEs
fs <- fscores(mod, full.scores.SE=TRUE)
head(fs)
empirical_rxx(fs)

## End(Not run)

```

MDIFF

*Compute multidimensional difficulty index***Description**

Returns a matrix containing the MDIFF values (Reckase, 2009). Only supported for items of class 'dich' and 'graded'.

Usage

```
MDIFF(x, which.items = NULL, group = NULL)
```

Arguments

x	an object of class 'SingleGroupClass', or an object of class 'MultipleGroup-Class' if a suitable group input were supplied
which.items	a vector indicating which items to select. If NULL is used (the default) then MDISC will be computed for all items
group	group argument to pass to extract.group function. Required when the input object is a multiple-group model

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Reckase, M. D. (2009). Multidimensional Item Response Theory. Springer.

See Also

[extract.group](#), [MDISC](#)

Examples

```
## Not run:

mod <- mirt(Science, 2)
MDIFF(mod)

mod <- mirt(expand.table(LSAT7), 2)
MDIFF(mod)

## End(Not run)
```

mdirt

Multidimensional discrete item response theory

Description

mdirt fits a variety of item response models with discrete latent variables. These include, but are not limited to, latent class analysis, multidimensional latent class models, multidimensional discrete latent class models, DINA/DINO models, grade of measurement models, C-RUM, and so on. If response models are not defined explicitly then customized models can be defined using the [createItem](#) function.

Usage

```
mdirt(
  data,
  model,
  customTheta = NULL,
  structure = NULL,
  item.Q = NULL,
  nruns = 1,
  method = "EM",
  covdata = NULL,
  formula = NULL,
  itemtype = "lca",
  optimizer = "nlminb",
```

```

    return_max = TRUE,
    group = NULL,
    GenRandomPars = FALSE,
    verbose = TRUE,
    pars = NULL,
    technical = list(),
    ...
)

```

Arguments

data	a matrix or data.frame that consists of numerically ordered data, organized in the form of integers, with missing data coded as NA
model	number of mutually exclusive classes to fit, or alternatively a more specific mirt.model definition (which reflects the so-called Q-matrix). Note that when using a mirt.model , the order with which the syntax factors/attributes are defined are associated with the columns in the customTheta input
customTheta	input passed to <code>technical = list(customTheta = ...)</code> , but is included directly in this function for convenience. This input is most interesting for discrete latent models because it allows customized patterns of latent classes (i.e., defines the possible combinations of the latent attribute profile). The default builds the pattern <code>customTheta = diag(model)</code> , which is the typical pattern for the traditional latent class analysis whereby class membership mutually distinct and exhaustive. See thetaComb for a quick method to generate a matrix with all possible combinations
structure	an R formula allowing the profile probability patterns (i.e., the structural component of the model) to be fitted according to a log-linear model. When NULL, all profile probabilities (except one) will be estimated. Use of this input requires that the customTheta input is supplied, and that the column names in this matrix match the names found within this formula
item.Q	a list of item-level Q-matrices indicating how the respective categories should be modeled by the underlying attributes. Each matrix must represent a $K_i \times A$ matrix, where K_i represents the number of categories for the i th item, and A is the number of attributes included in the Theta matrix; otherwise, a value of NULL will default to a matrix consisting of 1's for each $K_i \times A$ element except for the first row, which contains only 0's for proper identification. Incidentally, the first row of each matrix must contain only 0's so that the first category represents the reference category for identification
nruns	a numeric value indicating how many times the model should be fit to the data when using random starting values. If greater than 1, GenRandomPars is set to true by default
method	estimation method. Can be 'EM' or 'BL' (see mirt for more details)
covdata	a data.frame of data used for latent regression models
formula	an R formula (or list of formulas) indicating how the latent traits can be regressed using external covariates in covdata. If a named list of formulas is supplied (where the names correspond to the latent trait/attribute names in model) then

	specific regression effects can be estimated for each factor. Supplying a single formula will estimate the regression parameters for all latent variables by default
<code>itemtype</code>	a vector indicating the <code>itemtype</code> associated with each item. For discrete models this is limited to only 'lca' or items defined using a <code>createItem</code> definition
<code>optimizer</code>	optimizer used for the M-step, set to 'nllminb' by default. See <code>mirt</code> for more details
<code>return_max</code>	logical; when <code>nruns > 1</code> , return the model that has the most optimal maximum likelihood criteria? If FALSE, returns a list of all the estimated objects
<code>group</code>	a factor variable indicating group membership used for multiple group analyses
<code>GenRandomPars</code>	logical; use random starting values
<code>verbose</code>	logical; turn on messages to the R console
<code>pars</code>	used for modifying starting values; see <code>mirt</code> for details
<code>technical</code>	list of lower-level inputs. See <code>mirt</code> for details
<code>...</code>	additional arguments to be passed to the estimation engine. See <code>mirt</code> for more details and examples

Details

Posterior classification accuracy for each response pattern may be obtained via the `fscores` function. The `summary()` function will display the category probability values given the class membership, which can also be displayed graphically with `plot()`, while `coef()` displays the raw coefficient values (and their standard errors, if estimated). Finally, `anova()` is used to compare nested models, while `M2` and `itemfit` may be used for model fitting purposes.

'lca' model definition

The latent class IRT model with two latent classes has the form

$$P(x = k | \theta_1, \theta_2, a_1, a_2) = \frac{\exp(a_1\theta_1 + a_2\theta_2)}{\sum_j^K \exp(a_1\theta_1 + a_2\theta_2)}$$

where the θ values generally take on discrete points (such as 0 or 1). For proper identification, the first category slope parameters (a_1 and a_2) are never freely estimated. Alternatively, supplying a different grid of θ values will allow the estimation of similar models (multidimensional discrete models, grade of membership, etc.). See the examples below.

When the `item.Q` for is utilized, the above equation can be understood as

$$P(x = k | \theta_1, \theta_2, a_1, a_2) = \frac{\exp(a_1\theta_1 Q_{j1} + a_2\theta_2 Q_{j2})}{\sum_j^K \exp(a_1\theta_1 Q_{j1} + a_2\theta_2 Q_{j2})}$$

where by construction Q is a $K_i \times A$ matrix indicating whether the category should be modeled according to the latent class structure. For the standard latent class model, the Q -matrix has as many rows as categories, as many columns as the number of classes/attributes modeled, and consist of 0's in the first row and 1's elsewhere. This of course can be over-written by passing an alternative `item.Q` definition for each respective item.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29.

Proctor, C. H. (1970). A probabilistic formulation and statistical analysis for Guttman scaling. *Psychometrika*, 35, 73-78. doi:10.18637/jss.v048.i06

See Also

[thetaComb](#), [fcores](#), [mirt.model](#), [M2](#), [itemfit](#), [boot.mirt](#), [mirtCluster](#), [wald](#), [coef-method](#), [summary-method](#), [anova-method](#), [residuals-method](#)

Examples

```
# LSAT6 dataset
dat <- expand.table(LSAT6)

# fit with 2-3 latent classes
(mod2 <- mdirt(dat, 2))
## Not run:
(mod3 <- mdirt(dat, 3))
summary(mod2)
residuals(mod2)
residuals(mod2, type = 'exp')
anova(mod2, mod3)
M2(mod2)
itemfit(mod2)

# generate classification plots
plot(mod2)
plot(mod2, facet_items = FALSE)
plot(mod2, profile = TRUE)

# available for polytomous data
mod <- mdirt(Science, 2)
summary(mod)
plot(mod)
plot(mod, profile=TRUE)

# classification based on response patterns
fcores(mod2, full.scores = FALSE)

# classify individuals either with the largest posterior probability....
fs <- fcores(mod2)
head(fs)
classes <- 1:2
class_max <- classes[apply(apply(fs, 1, max) == fs, 1, which)]
table(class_max)
```

```

# ... or by probability sampling (i.e., plausible value draws)
class_prob <- apply(fs, 1, function(x) sample(1:2, 1, prob=x))
table(class_prob)

# plausible value imputations for stochastic classification in both classes
pvs <- fscores(mod2, plausible.draws=10)
tabs <- lapply(pvs, function(x) apply(x, 2, table))
tabs[[1]]

# fit with random starting points (run in parallel to save time)
if(interactive()) mirtCluster()
mod <- mdirt(dat, 2, nruns=10)

#-----
# Grade of measurement model

# define a custom Theta grid for including a 'fuzzy' class membership
(Theta <- matrix(c(1, 0, .5, .5, 0, 1), nrow=3, ncol=2, byrow=TRUE))
(mod_gom <- mdirt(dat, 2, customTheta = Theta))
summary(mod_gom)

#-----
# Multidimensional discrete latent class model

dat <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))

# define Theta grid for three latent classes
(Theta <- thetaComb(0:1, 3))
(mod_discrete <- mdirt(dat, 3, customTheta = Theta))
summary(mod_discrete)

# Located latent class model
model <- mirt.model('C1 = 1-32
                    C2 = 1-32
                    C3 = 1-32
                    CONSTRAIN = (1-32, a1), (1-32, a2), (1-32, a3)')
(mod_located <- mdirt(dat, model, customTheta = diag(3)))
summary(mod_located)

#-----
### DINA model example
# generate some suitable data for a two dimensional DINA application
#   (first columns are intercepts)
set.seed(1)
Theta <- expand.table(matrix(c(1,0,0,0,
                             1,1,0,0,
                             1,0,1,0,
                             1,1,1,1), 4, 4, byrow=TRUE),
  freq = c(200,200,100,500))
a <- matrix(c(rnorm(15, -1.5, .5), rlnorm(5, .2, .3), numeric(15), rlnorm(5, .2, .3),

```

```

        numeric(15), rlnorm(5, .2, .3)), 15, 4)

guess <- plogis(a[11:15,1]) # population guess
slip <- 1 - plogis(rowSums(a[11:15,])) # population slip

dat <- simdata(a, Theta=Theta, itemtype = 'lca')

# first column is the intercept, 2nd and 3rd are attributes
theta <- cbind(1, thetaComb(0:1, 2))
theta <- cbind(theta, theta[,2] * theta[,3]) #DINA interaction of main attributes
model <- mirt.model('Intercept = 1-15
                    A1 = 1-5
                    A2 = 6-10
                    A1A2 = 11-15')

# last 5 items are DINA (first 10 are unidimensional C-RUMs)
DINA <- mdirt(dat, model, customTheta = theta)
coef(DINA, simplify=TRUE)
summary(DINA)
M2(DINA) # fits well (as it should)

cfs <- coef(DINA, simplify=TRUE)$items[11:15,]
cbind(guess, estguess = plogis(cfs[,1]))
cbind(slip, estslip = 1 - plogis(rowSums(cfs)))

### DINO model example
theta <- cbind(1, thetaComb(0:1, 2))
# define theta matrix with negative interaction term
(theta <- cbind(theta, -theta[,2] * theta[,3]))

model <- mirt.model('Intercept = 1-15
                    A1 = 1-5, 11-15
                    A2 = 6-15
                    Yoshi = 11-15
                    CONSTRAIN = (11,a2,a3,a4), (12,a2,a3,a4), (13,a2,a3,a4),
                                (14,a2,a3,a4), (15,a2,a3,a4)')

# last five items are DINO's (first 10 are unidimensional C-RUMs)
DINO <- mdirt(dat, model, customTheta = theta)
coef(DINO, simplify=TRUE)
summary(DINO)
M2(DINO) #doesn't fit as well, because not the generating model

## C-RUM (analogous to MIRT model)
theta <- cbind(1, thetaComb(0:1, 2))
model <- mirt.model('Intercept = 1-15
                    A1 = 1-5, 11-15
                    A2 = 6-15')

CRUM <- mdirt(dat, model, customTheta = theta)
coef(CRUM, simplify=TRUE)
summary(CRUM)

```

```

# good fit, but over-saturated (main effects for items 11-15 can be set to 0)
M2(CRUM)

#-----
# multidimensional latent class model

dat <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))

# 5 latent classes within 2 different sets of items
model <- mirt.model('C1 = 1-16
  C2 = 1-16
  C3 = 1-16
  C4 = 1-16
  C5 = 1-16
  C6 = 17-32
  C7 = 17-32
  C8 = 17-32
  C9 = 17-32
  C10 = 17-32
  CONSTRAIN = (1-16, a1), (1-16, a2), (1-16, a3), (1-16, a4), (1-16, a5),
    (17-32, a6), (17-32, a7), (17-32, a8), (17-32, a9), (17-32, a10)')

theta <- diag(10) # defined explicitly. Otherwise, this profile is assumed
mod <- mdirt(dat, model, customTheta = theta)
coef(mod, simplify=TRUE)
summary(mod)

#-----
# multiple group with constrained group probabilities
dat <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
group <- rep(c('G1', 'G2'), each = nrow(SAT12)/2)
Theta <- diag(2)

# the latent class parameters are technically located in the (nitems + 1) location
model <- mirt.model('A1 = 1-32
  A2 = 1-32
  CONSTRAINB = (33, c1)')
mod <- mdirt(dat, model, group = group, customTheta = Theta)
coef(mod, simplify=TRUE)
summary(mod)

#-----
# Probabilistic Guttman Model (Proctor, 1970)

# example analysis can also be found in the sirt package (see ?prob.guttman)
data(data.read, package = 'sirt')
head(data.read)

Theta <- matrix(c(1,0,0,0,

```

```

      1,1,0,0,
      1,1,1,0,
      1,1,1,1), 4, byrow=TRUE)

model <- mirt.model("INTERCEPT = 1-12
                   C1 = 1,7,9,11
                   C2 = 2,5,8,10,12
                   C3 = 3,4,6")

mod <- mdirt(data.read, model, customTheta=Theta)
summary(mod)

M2(mod)
itemfit(mod)

## End(Not run)

```

MDISC

*Compute multidimensional discrimination index***Description**

Returns a vector containing the MDISC values for each item in the model input object (Reckase, 2009).

Usage

```
MDISC(x, group = NULL)
```

Arguments

x	an object of class 'SingleGroupClass', or an object of class 'MultipleGroup-Class' if a suitable group input were supplied
group	group argument to pass to extract.group function. Required when the input object is a multiple-group model

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Reckase, M. D. (2009). Multidimensional Item Response Theory. Springer.

See Also

[extract.group](#)

Examples

```
## Not run:

mod <- mirt(Science, 2)
MDISC(mod)

## End(Not run)
```

mirt	<i>Full-Information Item Factor Analysis (Multidimensional Item Response Theory)</i>
------	--

Description

mirt fits a maximum likelihood (or maximum a posteriori) factor analysis model to any mixture of dichotomous and polytomous data under the multidimensional item response theory paradigm using either Cai's (2010) Metropolis-Hastings Robbins-Monro (MHRM) algorithm, with an EM algorithm approach outlined by Bock and Aitkin (1981) using rectangular or quasi-Monte Carlo integration grids, or with the stochastic EM (i.e., the first two stages of the MH-RM algorithm). Unidimensional and multidimensional dominance/compensatory response models or unfolding/pairwise comparison models can be specified independently via the `itemtype` argument.

Usage

```
mirt(
  data,
  model = 1,
  itemtype = NULL,
  guess = 0,
  upper = 1,
  SE = FALSE,
  covdata = NULL,
  formula = NULL,
  itemdesign = NULL,
  item.formula = NULL,
  SE.type = "Oakes",
  method = "EM",
  optimizer = NULL,
  dentype = "Gaussian",
  pars = NULL,
  constrain = NULL,
  calcNull = FALSE,
```

```

draws = 5000,
survey.weights = NULL,
quadpts = NULL,
TOL = NULL,
gpcm_mats = list(),
grsm.block = NULL,
rsm.block = NULL,
monopoly.k = 1L,
key = NULL,
large = FALSE,
GenRandomPars = FALSE,
accelerate = "Ramsay",
verbose = TRUE,
solnp_args = list(),
nloptr_args = list(),
spline_args = list(),
control = list(),
technical = list(),
...
)

```

Arguments

- | | |
|----------|---|
| data | a matrix or data.frame that consists of numerically ordered data, organized in the form of integers, with missing data coded as NA (to convert from an ordered factor data.frame see data.matrix) |
| model | a string to be passed (or an object returned from) mirt.model , declaring how the IRT model is to be estimated (loadings, constraints, priors, etc). For exploratory IRT models, a single numeric value indicating the number of factors to extract is also supported. Default is 1, indicating that a unidimensional model will be fit unless otherwise specified |
| itemtype | <p>type of items to be modeled, declared as either a) a single value to be recycled for each item, b) a vector for each respective item, or c) if applicable, a matrix with columns equal to the number of items and rows equal to the number of latent classes. The NULL default assumes that the items follow a graded or 2PL structure, however they may be changed to the following:</p> <ul style="list-style-type: none"> • 'Rasch' - Rasch/partial credit model by constraining slopes to 1 and freely estimating the variance parameters (alternatively, can be specified by applying equality constraints to the slope parameters in 'gpcm' and '2PL'; Rasch, 1960) • '1PL', '2PL', '3PL', '3PLu', and '4PL' - 1-4 parameter logistic model, where 3PL estimates the lower asymptote only while 3PLu estimates the upper asymptote only (Lord and Novick, 1968; Lord, 1980). Note that specifying '1PL' will not automatically estimate the variance of the latent trait compared to the 'Rasch' type • '5PL' - 5 parameter logistic model to estimate asymmetric logistic response curves. Currently restricted to unidimensional models |

- 'CLL' - complementary log-log link model. Currently restricted to unidimensional models
- 'ULL' - unipolar log-logistic model (Lucke, 2015). Note the use of this itemtype will automatically use a log-normal distribution for the latent traits
- 'graded' - graded response model (Samejima, 1969)
- 'grsm' - graded ratings scale model in the classical IRT parameterization (restricted to unidimensional models; Muraki, 1992)
- 'gpcm' and 'gpcmIRT' - generalized partial credit model in the slope-intercept and classical parameterization. 'gpcmIRT' is restricted to unidimensional models. Note that optional scoring matrices for 'gpcm' are available with the `gpcm_mats` input (Muraki, 1992)
- 'rsm' - Rasch rating scale model using the 'gpcmIRT' structure (unidimensional only; Andrich, 1978)
- 'nominal' - nominal response model (Bock, 1972)
- 'ideal' - dichotomous ideal point model (Maydeu-Olivares, 2006)
- 'ggum' - generalized graded unfolding model (Roberts, Donoghue, & Laughlin, 2000) and its multidimensional extension
- 'hcm' and 'ghcm' - (generalized) hyperbolic cosine model (Andrich and Luo, 1993; Andrich, 1996) for dichotomous or ordered polytomous items (see Luo, 2001)
- 'alm' and 'galm' - (generalized) absolute logistic model (Luo and Andrich, 2005) for dichotomous or ordered polytomous items
- 'sslm' and 'gsslm' - (generalized) simple squared logistic model (Andrich, 1988) for dichotomous or ordered polytomous items (see Luo, 2001)
- 'paralla' and 'gparalla' - (generalized) parallellogram analysis model (Hojitink, 1990) for dichotomous or ordered polytomous items (see Luo, 2001)
- 'sequential' - multidimensional sequential response model (Tutz, 1990) in slope-intercept form
- 'Tutz' - same as the 'sequential' itemtype, except the slopes are fixed to 1 and the latent variance terms are freely estimated (similar to the 'Rasch' itemtype input)
- 'PC1PL', 'PC2PL', and 'PC3PL' - 1-3 parameter partially compensatory model. Note that constraining the slopes to be equal across items will also reduce the model to Embretson's (a.k.a. Whitely's) multicomponent model (1980), while for 'PC1PL' the slopes are fixed to 1 while the latent trait variance terms are estimated
- '2PLNRM', '3PLNRM', '3PLuNRM', and '4PLNRM' - 2-4 parameter nested logistic model, where 3PLNRM estimates the lower asymptote only while 3PLuNRM estimates the upper asymptote only (Suh and Bolt, 2010)
- 'spline' - spline response model with the `bs` (default) or the `ns` function (Winsberg, Thissen, and Wainer, 1984)
- 'monopoly' - monotonic polynomial model for unidimensional tests for dichotomous and polytomous response data (Falk and Cai, 2016)

Additionally, user defined item classes can also be defined using the `createItem` function

guess	fixed pseudo-guessing parameters. Can be entered as a single value to assign a global guessing parameter or may be entered as a numeric vector corresponding to each item
upper	fixed upper bound parameters for 4-PL model. Can be entered as a single value to assign a global guessing parameter or may be entered as a numeric vector corresponding to each item
SE	logical; estimate the standard errors by computing the parameter information matrix? See <code>SE.type</code> for the type of estimates available
covdata	a <code>data.frame</code> of data used for latent regression models
formula	an R formula (or list of formulas) indicating how the latent traits can be regressed using external covariates in <code>covdata</code> . If a named list of formulas is supplied (where the names correspond to the latent trait names in <code>model</code>) then specific regression effects can be estimated for each factor. Supplying a single formula will estimate the regression parameters for all latent traits by default
itemdesign	a <code>data.frame</code> with rows equal to the number of items and columns containing any item-design effects. If items should be included in the design structure (i.e., should be left in their canonical structure) then fewer rows can be used, however the rownames must be defined and matched with colnames in the data input. The item design matrix is constructed with the use of <code>item.formula</code> . Providing this input will fix the associated 'd' intercepts to 0, where applicable
item.formula	an R formula used to specify any intercept decomposition (e.g., the LLTM; Fischer, 1983). Note that only the right-hand side of the formula is required for compensatory models. For non-compensatory <code>itemtypes</code> (e.g., 'PC1PL') the formula must include the name of the latent trait in the left hand side of the expression to indicate which of the trait specification should have their intercepts decomposed (see MLTM; Embretson, 1984)
SE.type	type of estimation method to use for calculating the parameter information matrix for computing standard errors and <code>wald</code> tests. Can be: • 'Richardson', 'forward', or 'central' for the numerical Richardson, forward difference, and central difference evaluation of observed Hessian matrix • 'crossprod' and 'Louis' for standard error computations based on the variance of the Fisher scores as well as Louis' (1982) exact computation of the observed information matrix. Note that Louis' estimates can take a long time to obtain for large sample sizes and long tests • 'sandwich' for the sandwich covariance estimate based on the 'crossprod' and 'Oakes' estimates (see Chalmers, 2018, for details) • 'sandwich.Louis' for the sandwich covariance estimate based on the 'crossprod' and 'Louis' estimates • 'Oakes' for Oakes' (1999) method using a central difference approximation (see Chalmers, 2018, for details) • 'SEM' for the supplemented EM (disables the <code>accelerate</code> option automatically; EM only) • 'Fisher' for the expected information, 'complete' for information based on the complete-data Hessian used in EM algorithm

- 'MHRM' and 'FMHRM' for stochastic approximations of observed information matrix based on the Robbins-Monro filter or a fixed number of MHRM draws without the RM filter. These are the only options supported when `method = 'MHRM'`
- 'numerical' to obtain the numerical estimate from a call to `optim` when `method = 'BL'`

Note that both the 'SEM' method becomes very sensitive if the ML solution has not been reached with sufficient precision, and may be further sensitive if the history of the EM cycles is not stable/sufficient for convergence of the respective estimates. Increasing the number of iterations (increasing `NCYCLES` and decreasing `TOL`, see below) will help to improve the accuracy, and can be run in parallel if a `mirtCluster` object has been defined (this will be used for Oakes' method as well). Additionally, inspecting the symmetry of the ACOV matrix for convergence issues by passing `technical = list(symmetric = FALSE)` can be helpful to determine if a sufficient solution has been reached

method a character object specifying the estimation algorithm to be used. The default is 'EM', for the standard EM algorithm with fixed quadrature, 'QMCEM' for quasi-Monte Carlo EM estimation, or 'MCEM' for Monte Carlo EM estimation. The option 'MHRM' may also be passed to use the MH-RM algorithm, 'SEM' for the Stochastic EM algorithm (first two stages of the MH-RM stage using an optimizer other than a single Newton-Raphson iteration), and 'BL' for the Bock and Lieberman approach (generally not recommended for longer tests). The 'EM' is generally effective with 1-3 factors, but methods such as the 'QMCEM', 'MCEM', 'SEM', or 'MHRM' should be used when the dimensions are 3 or more. Note that when the optimizer is stochastic the associated `SE.type` is automatically changed to `SE.type = 'MHRM'` by default to avoid the use of quadrature

optimizer a character indicating which numerical optimizer to use. By default, the EM algorithm will use the 'BFGS' when there are no upper and lower bounds box-constraints and 'nllminb' when there are.

Other options include the Newton-Raphson ('NR'), which can be more efficient than the 'BFGS' but not as stable for more complex IRT models (such as the nominal or nested logit models) and the related 'NR1' which is also the Newton-Raphson but consists of only 1 update that has been coupled with RM Hessian (only applicable when the MH-RM algorithm is used). The MH-RM algorithm uses the 'NR1' by default, though currently the 'BFGS', 'L-BFGS-B', and 'NR' are also supported with this method (with fewer iterations by default) to emulate stochastic EM updates. As well, the 'Nelder-Mead' and 'SANN' estimators are available, but their routine use generally is not required or recommended.

Additionally, estimation subroutines from the `Rsolnp` and `nloptr` packages are available by passing the arguments 'solnp' and 'nloptr', respectively. This should be used in conjunction with the `solnp_args` and `nloptr_args` specified below. If equality constraints were specified in the model definition only the parameter with the lowest parnum in the `pars = 'values'` data.frame is used in the estimation vector passed to the objective function, and group hyper-parameters are omitted. Equality an inequality functions should be of the form `function(p, optim_args)`, where `optim_args` is a list of internally parameters that largely can be ignored when defining constraints (though use of `browser()` here may be helpful)

dentype	type of density form to use for the latent trait parameters. Current options include • 'Gaussian' (default) assumes a multivariate Gaussian distribution with an associated mean vector and variance-covariance matrix • 'empiricalhist' or 'EH' estimates latent distribution using an empirical histogram described by Bock and Aitkin (1981). Only applicable for uni-dimensional models estimated with the EM algorithm. For this option, the number of cycles, TOL, and quadpts are adjusted accommodate for less precision during estimation (namely: TOL = 3e-5, NCYCLES = 2000, quadpts = 121) • 'empiricalhist_Woods' or 'EHW' estimates latent distribution using an empirical histogram described by Bock and Aitkin (1981), with the same specifications as in dentype = 'empiricalhist', but with the extrapolation-interpolation method described by Woods (2007). NOTE: to improve stability in the presence of extreme response styles (i.e., all highest or lowest in each item) the technical option zeroExtreme = TRUE may be required to down-weight the contribution of these problematic patterns • 'Davidian-#' estimates semi-parametric Davidian curves described by Woods and Lin (2009), where the # placeholder represents the number of Davidian parameters to estimate (e.g., 'Davidian-6' will estimate 6 smoothing parameters). By default, the number of quadpts is increased to 121, and this method is only applicable for unidimensional models estimated with the EM algorithm Note that when itemtype = 'ULL' then a log-normal(0,1) density is used to support the unipolar scaling
pars	a data.frame with the structure of how the starting values, parameter numbers, estimation logical values, etc, are defined. The user may observe how the model defines the values by using pars = 'values', and this object can in turn be modified and input back into the estimation with pars = mymodifiedpars
constrain	a list of user declared equality constraints. To see how to define the parameters correctly use pars = 'values' initially to see how the parameters are labeled. To constrain parameters to be equal create a list with separate concatenated vectors signifying which parameters to constrain. For example, to set parameters 1 and 5 equal, and also set parameters 2, 6, and 10 equal use constrain = list(c(1, 5), c(2, 6, 10)). Constraints can also be specified using the mirt.model syntax (recommended)
calcNull	logical; calculate the Null model for additional fit statistics (e.g., TLI)? Only applicable if the data contains no NA's and the data is not overly sparse
draws	the number of Monte Carlo draws to estimate the log-likelihood for the MH-RM algorithm. Default is 5000
survey.weights	a optional numeric vector of survey weights to apply for each case in the data (EM estimation only). If not specified, all cases are weighted equally (the standard IRT approach). The sum of the survey.weights must equal the total sample size for proper weighting to be applied
quadpts	number of quadrature points per dimension (must be larger than 2). By default the number of quadrature uses the following scheme: switch(as.character(nfact),

	'1'=61, '2'=31, '3'=15, '4'=9, '5'=7, 3). However, if the method input is set to 'QMCEM' and this argument is left blank then the default number of quasi-Monte Carlo integration nodes will be set to 5000 in total
TOL	convergence threshold for EM or MH-RM; defaults are .0001 and .001. If SE.type = 'SEM' and this value is not specified, the default is set to 1e-5. To evaluate the model using only the starting values pass TOL = NaN, and to evaluate the starting values without the log-likelihood pass TOL = NA
gpcm_mats	a list of matrices specifying how the scoring coefficients in the (generalized) partial credit model should be constructed. If omitted, the standard gpcm format will be used (i.e., seq(0, k, by = 1) for each trait). This input should be used if traits should be scored different for each category (e.g., matrix(c(0:3, 1, 0, 0, 0), 4, 2) for a two-dimensional model where the first trait is scored like a gpcm, but the second trait is only positively indicated when the first category is selected). Can be used when itemtypes are 'gpcm' or 'Rasch', but only when the respective element in gpcm_mats is not NULL
grsm.block	an optional numeric vector indicating where the blocking should occur when using the grsm, NA represents items that do not belong to the grsm block (other items that may be estimated in the test data). For example, to specify two blocks of 3 with a 2PL item for the last item: grsm.block = c(rep(1, 3), rep(2, 3), NA). If NULL the all items are assumed to be within the same group and therefore have the same number of item categories
rsm.block	same as grsm.block, but for 'rsm' blocks
monopoly.k	a vector of values (or a single value to repeated for each item) which indicate the degree of the monotone polynomial fitted, where the monotone polynomial corresponds to monopoly.k * 2 + 1 (e.g., monopoly.k = 2 fits a 5th degree polynomial). Default is monopoly.k = 1, which fits a 3rd degree polynomial
key	a numeric vector of the response scoring key. Required when using nested logit item types, and must be the same length as the number of items used. Items that are not nested logit will ignore this vector, so use NA in item locations that are not applicable
large	a logical indicating whether unique response patterns should be obtained prior to performing the estimation so as to avoid repeating computations on identical patterns. The default TRUE provides the correct degrees of freedom for the model since all unique patterns are tallied (typically only affects goodness of fit statistics such as G2, but also will influence nested model comparison methods such as anova(mod1, mod2)), while FALSE will use the number of rows in data as a placeholder for the total degrees of freedom. As such, model objects should only be compared if all flags were set to TRUE or all were set to FALSE Alternatively, if the collapse table of frequencies is desired for the purpose of saving computations (i.e., only computing the collapsed frequencies for the data on-time) then a character vector can be passed with the argument large = 'return' to return a list of all the desired table information used by mirt. This list object can then be reused by passing it back into the large argument to avoid re-tallying the data again (again, useful when the dataset are very large and computing the tabulated data is computationally burdensome). This strategy is shown below:

Compute organized data e.g., `internaldat <- mirt(Science, 1, large = 'return')`

Pass the organized data to all estimation functions e.g., `mod <- mirt(Science, 1, large = internaldat)`

GenRandomPars	logical; generate random starting values prior to optimization instead of using the fixed internal starting values?
accelerate	a character vector indicating the type of acceleration to use. Default is 'Ramsay', but may also be 'squarem' for the SQUAREM procedure (specifically, the gSqS3 approach) described in Varadhan and Roldand (2008). To disable the acceleration, pass 'none'
verbose	logical; print observed- (EM) or complete-data (MHRM) log-likelihood after each iteration cycle? Default is TRUE
solnp_args	a list of arguments to be passed to the <code>solnp::solnp()</code> function for equality constraints, inequality constraints, etc
nloptr_args	a list of arguments to be passed to the <code>nloptr::nloptr()</code> function for equality constraints, inequality constraints, etc
spline_args	a named list of lists containing information to be passed to the <code>bs</code> (default) and <code>ns</code> for each spline itemtype. Each element must refer to the name of the itemtype with the spline, while the internal list names refer to the arguments which are passed. For example, if item 2 were called 'read2', and item 5 were called 'read5', both of which were of itemtype 'spline' but item 5 should use the <code>ns</code> form, then a modified list for each input might be of the form: <pre>spline_args = list(read2 = list(degree = 4), read5 = list(fun = 'ns', knots = c(-2, 2)))</pre> This code input changes the <code>bs()</code> splines function to have a degree = 4 input, while the second element changes to the <code>ns()</code> function with knots set a <code>c(-2, 2)</code>
control	a list passed to the respective optimizers (i.e., <code>optim()</code> , <code>nlminb()</code> , etc). Additional arguments have been included for the 'NR' optimizer: 'tol' for the convergence tolerance in the M-step (default is <code>TOL/1000</code>), while the default number of iterations for the Newton-Raphson optimizer is 50 (modified with the 'maxit' control input)
technical	a list containing lower level technical parameters for estimation. May be: NCYCLES maximum number of EM or MH-RM cycles; defaults are 500 and 2000 MAXQUAD maximum number of quadratures, which you can increase if you have more than 4GB or RAM on your PC; default 20000 theta_lim range of integration grid for each dimension; default is <code>c(-6, 6)</code>. Note that when itemtype = 'ULL' a log-normal distribution is used and the range is change to <code>c(.01, and 6^2)</code>, where the second term is the square of the <code>theta_lim</code> input instead set.seed seed number used during estimation. Default is 12345 SEtol standard error tolerance criteria for the S-EM and MHRM computation of the information matrix. Default is 1e-3 symmetric logical; force S-EM/Oakes information matrix estimates to be symmetric? Default is TRUE so that computation of standard errors are more

stable. Setting this to FALSE can help to detect solutions that have not reached the ML estimate

SEM_window ratio of values used to define the S-EM window based on the observed likelihood differences across EM iterations. The default is $c(0, 1 - SEtol)$, which provides nearly the very full S-EM window (i.e., nearly all EM cycles used). To use the a smaller SEM window change the window to to something like $c(.9, .999)$ to start at a point farther into the EM history

warn logical; include warning messages during estimation? Default is TRUE

message logical; include general messages during estimation? Default is TRUE

customK a numeric vector used to explicitly declare the number of response categories for each item. This should only be used when constructing mirt model for reasons other than parameter estimation (such as to obtain factor scores), and requires that the input data all have 0 as the lowest category. The format is the same as the `extract.mirt(mod, 'K')` slot in all converged models

customPriorFun a custom function used to determine the normalized density for integration in the EM algorithm. Must be of the form `function(Theta, Etable){...}`, and return a numeric vector with the same length as number of rows in Theta. The Etable input contains the aggregated table generated from the current E-step computations. For proper integration, the returned vector should sum to 1 (i.e., normalized). Note that if using the Etable it will be NULL on the first call, therefore the prior will have to deal with this issue accordingly

zeroExtreme logical; assign extreme response patterns a survey.weight of 0 (formally equivalent to removing these data vectors during estimation)? When `dentype = 'EHW'`, where Woods' extrapolation is utilized, this option may be required if the extrapolation causes expected densities to tend towards positive or negative infinity. The default is FALSE

customTheta a custom Theta grid, in matrix form, used for integration. If not defined, the grid is determined internally based on the number of quadpts

nconstrain same specification as the `constrain` list argument, however imposes a negative equality constraint instead (e.g., $a_{12} = -a_{21}$, which is specified as `nconstrain = list(c(12, 21))`). Note that each specification in the list must be of length 2, where the second element is taken to be -1 times the first element

delta the deviation term used in numerical estimates when computing the ACOV matrix with the 'forward' or 'central' numerical approaches, as well as Oakes' method with the Richardson extrapolation. Default is $1e-5$

parallel logical; use the parallel cluster defined by [mirtCluster](#)? Default is TRUE

storeEMhistory logical; store the iteration history when using the EM algorithm? Default is FALSE. When TRUE, use [extract.mirt](#) to extract

internal_constraints logical; include the internal constraints when using certain IRT models (e.g., 'grsm' itemtype). Disable this if you want to use special optimizers such as the `solnp`. Default is TRUE

gain a vector of two values specifying the numerator and exponent values for the RM gain function $(val1/cycle)^{val2}$. Default is $c(0.10, 0.75)$

BURNIN number of burn in cycles (stage 1) in MH-RM; default is 150

SEMCYCLES number of SEM cycles (stage 2) in MH-RM; default is 100

MHDRAWS number of Metropolis-Hasting draws to use in the MH-RM at each iteration; default is 5

MHcand a vector of values used to tune the MH sampler. Larger values will cause the acceptance ratio to decrease. One value is required for each group in unconditional item factor analysis (`mixedmirt()` requires additional values for random effect). If null, these values are determined internally, attempting to tune the acceptance of the draws to be between .1 and .4

MHRM_SE_draws number of fixed draws to use when `SE=TRUE` and `SE.type = 'FMHRM'` and the maximum number of draws when `SE.type = 'MHRM'`. Default is 2000

MCEM_draws a function used to determine the number of quadrature points to draw for the 'MCEM' method. Must include one argument which indicates the iteration number of the EM cycle. Default is `function(cycles) 500 + (cycles - 1)*2`, which starts the number of draws at 500 and increases by 2 after each full EM iteration

info_if_converged logical; compute the information matrix when using the MH-RM algorithm only if the model converged within a suitable number of iterations? Default is TRUE

logLik_if_converged logical; compute the observed log-likelihood when using the MH-RM algorithm only if the model converged within a suitable number of iterations? Default is TRUE

keep_vcov_PD logical; attempt to keep the variance-covariance matrix of the latent traits positive definite during estimation in the EM algorithm? This generally improves the convergence properties when the traits are highly correlated. Default is TRUE

... additional arguments to be passed

Details

Models containing 'explanatory' person or item level predictors can only be included by using the `mixedmirt` function, though latent regression models can be fit using the formula input in this function. Tests that form a two-tier or bi-factor structure should be estimated with the `bfactor` function, which uses a dimension reduction EM algorithm for modeling item parcels. Multiple group analyses (useful for DIF and DTF testing) are also available using the `multipleGroup` function.

Value

function returns an object of class `SingleGroupClass` (`SingleGroupClass-class`)

Confirmatory and Exploratory IRT

Specification of the confirmatory item factor analysis model follows many of the rules in the structural equation modeling framework for confirmatory factor analysis. The variances of the latent factors are automatically fixed to 1 to help facilitate model identification. All parameters may be

fixed to constant values or set equal to other parameters using the appropriate declarations. Confirmatory models may also contain 'explanatory' person or item level predictors, though including predictors is currently limited to the `mixedmirt` function.

When specifying a single number greater than 1 as the model input to mirt an exploratory IRT model will be estimated. Rotation and target matrix options are available if they are passed to generic functions such as `summary-method` and `fscores`. Factor means and variances are fixed to ensure proper identification.

If the model is an exploratory item factor analysis estimation will begin by computing a matrix of quasi-polychoric correlations. A factor analysis with `nfact` is then extracted and item parameters are estimated by $a_{ij} = f_{ij}/u_j$, where f_{ij} is the factor loading for the j th item on the i th factor, and u_j is the square root of the factor uniqueness, $\sqrt{1 - h_j^2}$. The initial intercept parameters are determined by calculating the inverse normal of the item facility (i.e., item easiness), q_j , to obtain $d_j = q_j/u_j$. A similar implementation is also used for obtaining initial values for polytomous items.

A note on upper and lower bound parameters

Internally the g and u parameters are transformed using a logit transformation ($\log(x/(1 - x))$), and can be reversed by using $1/(1 + \exp(-x))$ following convergence. This also applies when computing confidence intervals for these parameters, and is done so automatically if `coef(mod, rawug = FALSE)`.

As such, when applying prior distributions to these parameters it is recommended to use a prior that ranges from negative infinity to positive infinity, such as the normally distributed prior via the 'norm' input (see `mirt.model`).

Convergence for quadrature methods

Unrestricted full-information factor analysis is known to have problems with convergence, and some items may need to be constrained or removed entirely to allow for an acceptable solution. As a general rule dichotomous items with means greater than .95, or items that are only .05 greater than the guessing parameter, should be considered for removal from the analysis or treated with prior parameter distributions. The same type of reasoning is applicable when including upper bound parameters as well. For polytomous items, if categories are rarely endorsed then this will cause similar issues. Also, increasing the number of quadrature points per dimension, or using the quasi-Monte Carlo integration method, may help to stabilize the estimation process in higher dimensions. Finally, solutions that are not well defined also will have difficulty converging, and can indicate that the model has been misspecified (e.g., extracting too many dimensions).

Convergence for MH-RM method

For the MH-RM algorithm, when the number of iterations grows very high (e.g., greater than 1500) or when `Max Change = .2500` values are repeatedly printed to the console too often (indicating that the parameters were being constrained since they are naturally moving in steps greater than 0.25) then the model may either be ill defined or have a very flat likelihood surface, and genuine maximum-likelihood parameter estimates may be difficult to find. Standard errors are computed following the model convergence by passing `SE = TRUE`, to perform an additional MH-RM stage but treating the maximum-likelihood estimates as fixed points.

Additional helper functions

Additional functions are available in the package which can be useful pre- and post-estimation. These are:

- `mirt.model` Define the IRT model specification use special syntax. Useful for defining between and within group parameter constraints, prior parameter distributions, and specifying the slope coefficients for each factor
- `coef-method` Extract raw coefficients from the model, along with their standard errors and confidence intervals
- `summary-method` Extract standardized loadings from model. Accepts a rotate argument for exploratory item response model
- `anova-method` Compare nested models using likelihood ratio statistics as well as information criteria such as the AIC and BIC
- `residuals-method` Compute pairwise residuals between each item using methods such as the LD statistic (Chen & Thissen, 1997), as well as response pattern residuals
- `plot-method` Plot various types of test level plots including the test score and information functions and more
- `itemplot` Plot various types of item level plots, including the score, standard error, and information functions, and more
- `createItem` Create a customized itemtype that does not currently exist in the package
- `imputeMissing` Impute missing data given some computed Theta matrix
- `fscores` Find predicted scores for the latent traits using estimation methods such as EAP, MAP, ML, WLE, and EAPsum
- `wald` Compute Wald statistics follow the convergence of a model with a suitable information matrix
- `M2` Limited information goodness of fit test statistic based to determine how well the model fits the data
- `itemfit and personfit` Goodness of fit statistics at the item and person levels, such as the S-X2, infit, outfit, and more
- `boot.mirt` Compute estimated parameter confidence intervals via the bootstrap methods
- `mirtCluster` Define a cluster for the package functions to use for capitalizing on multi-core architecture to utilize available CPUs when possible. Will help to decrease estimation times for tasks that can be run in parallel

IRT Models

The parameter labels use the follow convention, here using two factors and K as the total number of categories (using k for specific category instances).

Rasch Only one intercept estimated, and the latent variance of θ is freely estimated. If the data have more than two categories then a partial credit model is used instead (see 'gpcm' below).

$$P(x = 1|\theta, d) = \frac{1}{1 + \exp(-(\theta + d))}$$

1-4PL Depending on the model u may be equal to 1 (e.g., 3PL), g may be equal to 0 (e.g., 2PL), or the a may be fixed to 1 (e.g., 1PL).

$$P(x = 1|\theta, \psi) = g + \frac{(u - g)}{1 + \exp(-(a_1 * \theta_1 + a_2 * \theta_2 + d))}$$

5PL Currently restricted to unidimensional models

$$P(x = 1|\theta, \psi) = g + \frac{(u - g)}{1 + \exp(-(a_1 * \theta_1 + d))^S}$$

where S allows for asymmetry in the response function and is transformation constrained to be greater than 0 (i.e., $\log(S)$ is estimated rather than S)

CLL Complementary log-log model (see Shim, Bonifay, and Wiedermann, 2022)

$$P(x = 1|\theta, b) = 1 - \exp(-\exp(\theta - b))$$

Currently restricted to unidimensional dichotomous data.

graded The graded model consists of sequential 2PL models,

$$P(x = k|\theta, \psi) = P(x \geq k|\theta, \phi) - P(x \geq k + 1|\theta, \phi)$$

Note that $P(x \geq 1|\theta, \phi) = 1$ while $P(x \geq K + 1|\theta, \phi) = 0$

ULL The unipolar log-logistic model (ULL; Lucke, 2015) is defined the same as the graded response model, however

$$P(x \leq k|\theta, \psi) = \frac{\lambda_k \theta^\eta}{1 + \lambda_k \theta^\eta}$$

. Internally the λ parameters are exponentiated to keep them positive, and should therefore the reported estimates should be interpreted in log units

grsm A more constrained version of the graded model where graded spacing is equal across item blocks and only adjusted by a single 'difficulty' parameter (c) while the latent variance of θ is freely estimated (see Muraki, 1990 for this exact form). This is restricted to unidimensional models only.

gpcm/nominal For the gpcm the d values are treated as fixed and ordered values from $0 : (K - 1)$ (in the nominal model d_0 is also set to 0). Additionally, for identification in the nominal model $ak_0 = 0, ak_{(K-1)} = (K - 1)$.

$$P(x = k|\theta, \psi) = \frac{\exp(ak_{k-1} * (a_1 * \theta_1 + a_2 * \theta_2) + d_{k-1})}{\sum_{k=1}^K \exp(ak_{k-1} * (a_1 * \theta_1 + a_2 * \theta_2) + d_{k-1})}$$

For the partial credit model (when `itemtype = 'Rasch'`; unidimensional only) the above model is further constrained so that $ak = (0, 1, \dots, K - 1)$, $a_1 = 1$, and the latent variance of θ_1 is freely estimated. Alternatively, the partial credit model can be obtained by containing all the slope parameters in the gpcms to be equal. More specific scoring function may be included by passing a suitable list or matrices to the `gpcm_mats` input argument.

In the nominal model this parametrization helps to identify the empirical ordering of the categories by inspecting the ak values. Larger values indicate that the item category is more positively related to the latent trait(s) being measured. For instance, if an item was truly ordinal (such as a Likert scale), and had 4 response categories, we would expect to see

$ak_0 < ak_1 < ak_2 < ak_3$ following estimation. If on the other hand $ak_0 > ak_1$ then it would appear that the second category is less related to the trait than the first, and therefore the second category should be understood as the 'lowest score'.

NOTE: The nominal model can become numerical unstable if poor choices for the high and low values are chosen, resulting in ak values greater than $\text{abs}(10)$ or more. It is recommended to choose high and low anchors that cause the estimated parameters to fall between 0 and $K - 1$ either by theoretical means or by re-estimating the model with better values following convergence.

gpcmIRT and rsm The gpcmIRT model is the classical generalized partial credit model for unidimensional response data. It will obtain the same fit as the gpcm presented above, however the parameterization allows for the Rasch/generalized rating scale model as a special case.

E.g., for a $K = 4$ category response model,

$$\begin{aligned} P(x = 0|\theta, \psi) &= \exp(0)/G \\ P(x = 1|\theta, \psi) &= \exp(a(\theta - b1) + c)/G \\ P(x = 2|\theta, \psi) &= \exp(a(2\theta - b1 - b2) + 2c)/G \\ P(x = 3|\theta, \psi) &= \exp(a(3\theta - b1 - b2 - b3) + 3c)/G \end{aligned}$$

where

$$G = \exp(0) + \exp(a(\theta - b1) + c) + \exp(a(2\theta - b1 - b2) + 2c) + \exp(a(3\theta - b1 - b2 - b3) + 3c)$$

Here a is the slope parameter, the b parameters are the threshold values for each adjacent category, and c is the so-called difficulty parameter when a rating scale model is fitted (otherwise, $c = 0$ and it drops out of the computations).

The gpcmIRT can be constrained to the partial credit IRT model by either constraining all the slopes to be equal, or setting the slopes to 1 and freeing the latent variance parameter.

Finally, the rsm is a more constrained version of the (generalized) partial credit model where the spacing is equal across item blocks and only adjusted by a single 'difficulty' parameter (c). Note that this is analogous to the relationship between the graded model and the grsm (with an additional constraint regarding the fixed discrimination parameters).

sequential/Tutz The multidimensional sequential response model has the form

$$P(x = k|\theta, \psi) = \prod (1 - F(a_1\theta_1 + a_2\theta_2 + d_{sk}))F(a_1\theta_1 + a_2\theta_2 + d_{jk})$$

where $F(\cdot)$ is the cumulative logistic function. The Tutz variant of this model (Tutz, 1990) (via `itemtype = 'Tutz'`) assumes that the slope terms are all equal to 1 and the latent variance terms are estimated (i.e., is a Rasch variant).

ideal The ideal point model has the form, with the upper bound constraint on d set to 0:

$$P(x = 1|\theta, \psi) = \exp(-0.5 * (a_1 * \theta_1 + a_2 * \theta_2 + d)^2)$$

partcomp Partially compensatory models consist of the product of 2PL probability curves.

$$P(x = 1|\theta, \psi) = g + (1 - g) \left(\frac{1}{1 + \exp(-(a_1 * \theta_1 + d_1))} \right)^{c_1} * \left(\frac{1}{1 + \exp(-(a_2 * \theta_2 + d_2))} \right)^{c_2}$$

where c_1 and c_2 are binary indicator variables reflecting whether the item should include the select compensatory component (1) or not (0). Note that constraining the slopes to be equal across items will reduce the model to Embretson's (Whitely's) multicomponent model (1980).

2-4PLNRM Nested logistic curves for modeling distractor items. Requires a scoring key. The model is broken into two components for the probability of endorsement. For successful endorsement the probability trace is the 1-4PL model, while for unsuccessful endorsement:

$$P(x = 0|\theta, \psi) = (1 - P_{1-4PL}(x = 1|\theta, \psi)) * P_{nominal}(x = k|\theta, \psi)$$

which is the product of the complement of the dichotomous trace line with the nominal response model. In the nominal model, the slope parameters defined above are constrained to be 1's, while the last value of the ak is freely estimated.

ggum The (multidimensional) generalized graded unfolding model is a class of ideal point models useful for ordinal response data. The form is

$$P(z = k|\theta, \psi) = \frac{\exp \left[\left(z \sqrt{\sum_{d=1}^D a_{id}^2 (\theta_{jd} - b_{id})^2} \right) + \sum_{k=0}^z \psi_{ik} \right] + \exp \left[\left((M - z) \sqrt{\sum_{d=1}^D a_{id}^2 (\theta_{jd} - b_{id})^2} \right) + \sum_{k=0}^{M-z} \psi_{ik} \right]}{\sum_{w=0}^C \left(\exp \left[\left(w \sqrt{\sum_{d=1}^D a_{id}^2 (\theta_{jd} - b_{id})^2} \right) + \sum_{k=0}^w \psi_{ik} \right] + \exp \left[\left((M - w) \sqrt{\sum_{d=1}^D a_{id}^2 (\theta_{jd} - b_{id})^2} \right) + \sum_{k=0}^{M-w} \psi_{ik} \right] \right)}$$

where θ_{jd} is the location of the j th individual on the d th dimension, b_{id} is the difficulty location of the i th item on the d th dimension, a_{id} is the discrimination of the j th individual on the d th dimension (where the discrimination values are constrained to be positive), ψ_{ik} is the k th subjective response category threshold for the i th item, assumed to be symmetric about the item and constant across dimensions, where $\psi_{ik} = \sum_{d=1}^D a_{id} t_{ik}$ $z = 1, 2, \dots, C$ (where C is the number of categories minus 1), and $M = 2C + 1$.

(g)hcm, (g)alm, (g)sslm, and (g)paralla Following Luo (2001), this family of response models can be characterized under the same ordinal response functioning structure, differing only in their linking functions ($\psi(x)$). For example, for a two-dimensional model the equation used is

$$p_k = \frac{\psi(\rho_k)}{\psi(\rho_k) + \psi(a_1\theta_1 + a_2\theta_2 + d)}$$

which is expressed in slope-intercept form to accommodate multidimensionality. The "generalized" versions of this family estimate the slope and rho parameters, which allow each item to differ in the steepness of the unfolding model functions; otherwise, slopes are fixed to the value of 1, though the rho parameters must be set to 0 manually. For ordered polytomous items the response function follows the Guttman-scaling logic

$$P_1 = q_1 \cdot q_2 \cdot q_3 \cdots q_k$$

$$P_2 = p_1 \cdot q_2 \cdot q_3 \cdots q_k$$

$$P_3 = p_1 \cdot p_2 \cdot q_3 \cdots q_k$$

...

$$P_K = p_1 \cdot p_2 \cdot p_3 \cdots p_k$$

Note that for estimation purposes the rho parameters are expressed in log units so that they remain positive during estimation. Hence, the parameterization used herein is

$$p_k = \frac{\psi(\exp(\rho_k^*))}{\psi(\exp(\rho_k^*)) + \psi(a_1\theta_1 + a_2\theta_2 + d)}$$

where ρ_k^* is in natural log units. Currently supported models in this family are the:

- (generalized) hyperbolic cosine model ($\psi(x) = \cosh(x)$),
- (generalized) absolute logistic model ($\psi(x) = \exp(|x|)$),
- (generalized) simple squared logistic model ($\psi(x) = \exp(x^2)$), and
- (generalized) parallellogram analysis model ($\psi(x) = x^2$), respectively.

all of which are available for dichotomous and ordered polytomous response option items.

spline Spline response models attempt to model the response curves uses non-linear and potentially non-monotonic patterns. The form is

$$P(x = 1|\theta, \eta) = \frac{1}{1 + \exp(-(\eta_1 * X_1 + \eta_2 * X_2 + \dots + \eta_n * X_n))}$$

where the X_n are from the spline design matrix X organized from the grid of θ values. B-splines with a natural or polynomial basis are supported, and the intercept input is set to TRUE by default.

monopoly Monotone polynomial model for polytomous response data of the form

$$P(x = k|\theta, \psi) = \frac{\exp(\sum_1^k (m^*(\psi) + \xi_{c-1}))}{\sum_1^C \exp(\sum_1^K (m^*(\psi) + \xi_{c-1}))}$$

where $m^*(\psi)$ is the monotone polynomial function without the intercept.

HTML help files, exercises, and examples

To access examples, vignettes, and exercise files that have been generated with knitr please visit <https://github.com/philchalmers/mirt/wiki>.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Andrich, D. (1978). A rating scale formulation for ordered response categories. *Psychometrika*, 43, 561-573.
- Andrich, D. (1996). Hyperbolic cosine latent trait models for unfolding direct-responses and pair-wise preferences. *Applied Psychological Measurement*, 20, 269-290.
- Andrich, D., and Luo, G. (1993). A hyperbolic cosine latent trait model for unfolding dichotomous single- stimulus responses. *Applied Psychological Measurement*, 17, 253-276.
- Andrich, D. (1988). The application of an unfolding model of the PIRT type to the measurement of attitude. *Applied Psychological Measurement*, 12, 33-51.
- Bock, R. D., & Aitkin, M. (1981). Marginal maximum likelihood estimation of item parameters: Application of an EM algorithm. *Psychometrika*, 46(4), 443-459.
- Bock, R. D., Gibbons, R., & Muraki, E. (1988). Full-Information Item Factor Analysis. *Applied Psychological Measurement*, 12(3), 261-280.
- Bock, R. D. & Lieberman, M. (1970). Fitting a response model for n dichotomously scored items. *Psychometrika*, 35, 179-197.

- Cai, L. (2010a). High-Dimensional exploratory item factor analysis by a Metropolis-Hastings Robbins-Monro algorithm. *Psychometrika*, 75, 33-57.
- Cai, L. (2010b). Metropolis-Hastings Robbins-Monro algorithm for confirmatory item factor analysis. *Journal of Educational and Behavioral Statistics*, 35, 307-335.
- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Chalmers, R. P. (2015). Extended Mixed-Effects Item Response Models with the MH-RM Algorithm. *Journal of Educational Measurement*, 52, 200-222. doi:10.1111/jedm.12072
- Chalmers, R. P. (2018). Numerical Approximation of the Observed Information Matrix with Oakes' Identity. *British Journal of Mathematical and Statistical Psychology* DOI: 10.1111/bmsp.12127
- Chalmers, R., P. & Flora, D. (2014). Maximum-likelihood Estimation of Noncompensatory IRT Models with the MH-RM Algorithm. *Applied Psychological Measurement*, 38, 339-358. doi:10.1177/0146621614520958
- Chen, W. H. & Thissen, D. (1997). Local dependence indices for item pairs using item response theory. *Journal of Educational and Behavioral Statistics*, 22, 265-289.
- Embretson, S. E. (1984). A general latent trait model for response processes. *Psychometrika*, 49, 175-186.
- Falk, C. F. & Cai, L. (2016). Maximum Marginal Likelihood Estimation of a Monotonic Polynomial Generalized Partial Credit Model with Applications to Multiple Group Analysis. *Psychometrika*, 81, 434-460.
- Fischer, G. H. (1983). Logistic latent trait models with linear constraints. *Psychometrika*, 48, 3-26.
- Hoijtink H. (1990). PARELLA: Measurement of latent traits by proximity items. The Netherlands: University of Groningen.
- Lord, F. M. & Novick, M. R. (1968). Statistical theory of mental test scores. Addison-Wesley.
- Lucke, J. F. (2015). Unipolar item response models. In S. P. Reise & D. A. Revicki (Eds.), *Handbook of item response theory modeling: Applications to typical performance assessment* (pp. 272-284). New York, NY: Routledge/Taylor & Francis Group.
- Luo G. (2001). A class of probabilistic unfolding models for polytomous responses. *Journal of Mathematical Psychology*. 45(2):224-248. 10.1006/jmps.2000.1310
- Luo G, and Andrich D. (2005). Information functions for the general dichotomous unfolding model. In: Alagumalai S, Curtis D.D., & Hungi N., editor. *Applied Rasch Measurement: A Book of Exemplars*: Dordrecht, The Netherlands: Springer.
- Maydeu-Olivares, A., Hernandez, A. & McDonald, R. P. (2006). A Multidimensional Ideal Point Item Response Theory Model for Binary Data. *Multivariate Behavioral Research*, 41, 445-471.
- Muraki, E. (1990). Fitting a polytomous item response model to Likert-type data. *Applied Psychological Measurement*, 14, 59-71.
- Muraki, E. (1992). A generalized partial credit model: Application of an EM algorithm. *Applied Psychological Measurement*, 16, 159-176.
- Muraki, E. & Carlson, E. B. (1995). Full-information factor analysis for polytomous item responses. *Applied Psychological Measurement*, 19, 73-90.
- Ramsay, J. O. (1975). Solving implicit equations in psychometric data analysis. *Psychometrika*, 40, 337-360.

- Rasch, G. (1960). Probabilistic models for some intelligence and attainment tests. *Danish Institute for Educational Research*.
- Roberts, J. S., Donoghue, J. R., & Laughlin, J. E. (2000). A General Item Response Theory Model for Unfolding Unidimensional Polytomous Responses. *Applied Psychological Measurement*, 24, 3-32.
- Samejima, F. (1969). Estimation of latent ability using a response pattern of graded scores. *Psychometrika Monographs*, 34.
- Shim, H., Bonifay, W., & Wiedermann, W. (2022). Parsimonious asymmetric item response theory modeling with the complementary log-log link. *Behavior Research Methods*, 55, 200-219.
- Suh, Y. & Bolt, D. (2010). Nested logit models for multiple-choice item response data. *Psychometrika*, 75, 454-473.
- Sympson, J. B. (1977). A model for testing with multidimensional items. Proceedings of the 1977 Computerized Adaptive Testing Conference.
- Thissen, D. (1982). Marginal maximum likelihood estimation for the one-parameter logistic model. *Psychometrika*, 47, 175-186.
- Tutz, G. (1990). Sequential item response models with ordered response. *British Journal of Mathematical and Statistical Psychology*, 43, 39-55.
- Varadhan, R. & Roland, C. (2008). Simple and Globally Convergent Methods for Accelerating the Convergence of Any EM Algorithm. *Scandinavian Journal of Statistics*, 35, 335-353.
- Whitely, S. E. (1980). Multicomponent latent trait models for ability tests. *Psychometrika*, 45(4), 470-494.
- Wood, R., Wilson, D. T., Gibbons, R. D., Schilling, S. G., Muraki, E., & Bock, R. D. (2003). *TESTFACT 4 for Windows: Test Scoring, Item Statistics, and Full-information Item Factor Analysis* [Computer software]. Lincolnwood, IL: Scientific Software International.
- Woods, C. M., and Lin, N. (2009). Item Response Theory With Estimation of the Latent Density Using Davidian Curves. *Applied Psychological Measurement*, 33(2), 102-117.

See Also

[bfactor](#), [multipleGroup](#), [mixedmirt](#), [expand.table](#), [key2binary](#), [mod2values](#), [extract.item](#), [iteminfo](#), [testinfo](#), [probtrace](#), [simdata](#), [averageMI](#), [fixef](#), [extract.mirt](#), [itemstats](#)

Examples

```
# load LSAT section 7 data and compute 1 and 2 factor models
data <- expand.table(LSAT7)
itemstats(data)

(mod1 <- mirt(data, 1))
coef(mod1)
summary(mod1)
plot(mod1)
plot(mod1, type = 'trace')

## Not run:
(mod2 <- mirt(data, 1, SE = TRUE)) #standard errors via the Oakes method
```

```

(mod2 <- mirt(data, 1, SE = TRUE, SE.type = 'SEM')) #standard errors with SEM method
coef(mod2)
(mod3 <- mirt(data, 1, SE = TRUE, SE.type = 'Richardson')) #with numerical Richardson method
residuals(mod1)
plot(mod1) #test score function
plot(mod1, type = 'trace') #trace lines
plot(mod2, type = 'info') #test information
plot(mod2, MI=200) #expected total score with 95% confidence intervals

# estimated 3PL model for item 5 only
(mod1.3PL <- mirt(data, 1, itemtype = c('2PL', '2PL', '2PL', '2PL', '3PL')))
coef(mod1.3PL)

# internally g and u pars are stored as logits, so usually a good idea to include normal prior
# to help stabilize the parameters. For a value around .182 use a mean
# of -1.5 (since  $1 / (1 + \exp(-(-1.5))) = .182$ )
model <- 'F = 1-5
PRIOR = (5, g, norm, -1.5, 3)'
mod1.3PL.norm <- mirt(data, model, itemtype = c('2PL', '2PL', '2PL', '2PL', '3PL'))
coef(mod1.3PL.norm)
#limited information fit statistics
M2(mod1.3PL.norm)

# unidimensional ideal point model
idealpt <- mirt(data, 1, itemtype = 'ideal')
plot(idealpt, type = 'trace', facet_items = TRUE)
plot(idealpt, type = 'trace', facet_items = FALSE)

# two factors (exploratory)
mod2 <- mirt(data, 2)
coef(mod2)
summary(mod2, rotate = 'oblimin') #oblimin rotation
residuals(mod2)
plot(mod2)
plot(mod2, rotate = 'oblimin')

anova(mod1, mod2) #compare the two models
scoresfull <- fscores(mod2) #factor scores for each response pattern
head(scoresfull)
scorestable <- fscores(mod2, full.scores = FALSE) #save factor score table
head(scorestable)

# confirmatory (as an example, model is not identified since you need 3 items per factor)
# Two ways to define a confirmatory model: with mirt.model, or with a string

# these model definitions are equivalent
cmodel <- mirt.model('
  F1 = 1,4,5
  F2 = 2,3')
cmodel2 <- 'F1 = 1,4,5
  F2 = 2,3'

cmod <- mirt(data, cmodel)

```

```

# cmod <- mirt(data, cmodel2) # same as above
coef(cmod)
anova(cmod, mod2)
# check if identified by computing information matrix
(cmod <- mirt(data, cmodel, SE = TRUE))

#####
# data from the 'ltm' package in numeric format
itemstats(Science)

pmod1 <- mirt(Science, 1)
plot(pmod1)
plot(pmod1, type = 'trace')
plot(pmod1, type = 'itemscore')
summary(pmod1)

# Constrain all slopes to be equal with the constrain = list() input or mirt.model() syntax
# first obtain parameter index
values <- mirt(Science, 1, pars = 'values')
values #note that slopes are numbered 1,5,9,13, or index with values$parnum[values$name == 'a1']
(pmod1_equalslopes <- mirt(Science, 1, constrain = list(c(1,5,9,13))))
coef(pmod1_equalslopes)

# using mirt.model syntax, constrain all item slopes to be equal
model <- 'F = 1-4
        CONSTRAIN = (1-4, a1)'
(pmod1_equalslopes <- mirt(Science, model))
coef(pmod1_equalslopes)

coef(pmod1_equalslopes)
anova(pmod1_equalslopes, pmod1) #significantly worse fit with almost all criteria

pmod2 <- mirt(Science, 2)
summary(pmod2)
plot(pmod2, rotate = 'oblimin')
itemplot(pmod2, 1, rotate = 'oblimin')
anova(pmod1, pmod2)

# unidimensional fit with a generalized partial credit and nominal model
(gpcmod <- mirt(Science, 1, 'gpcm'))
coef(gpcmod)

# for the nominal model the lowest and highest categories are assumed to be the
# theoretically lowest and highest categories that related to the latent trait(s)
(nomod <- mirt(Science, 1, 'nominal'))
coef(nomod) #ordering of ak values suggest that the items are indeed ordinal
anova(gpcmod, nomod)
itemplot(nomod, 3)

# generalized graded unfolding model
(ggum <- mirt(Science, 1, 'ggum'))
coef(ggum, simplify=TRUE)
plot(ggum)

```

```

plot(ggum, type = 'trace')
plot(ggum, type = 'itemscore')

# monotonic polyomial models
(monopoly <- mirt(Science, 1, 'monopoly'))
coef(monopoly, simplify=TRUE)
plot(monopoly)
plot(monopoly, type = 'trace')
plot(monopoly, type = 'itemscore')

# unipolar IRT model
unimod <- mirt(Science, itemtype = 'ULL')
coef(unimod, simplify=TRUE)
plot(unimod)
plot(unimod, type = 'trace')
itemplot(unimod, 1)

# following use the correct log-normal density for latent trait
itemfit(unimod)
M2(unimod, type = 'C2')
fs <- fscores(unimod)
hist(fs, 20)
fscores(unimod, method = 'EAPsum', full.scores = FALSE)

## example applying survey weights.
# weight the first half of the cases to be more representative of population
survey.weights <- c(rep(2, nrow(Science)/2), rep(1, nrow(Science)/2))
survey.weights <- survey.weights/sum(survey.weights) * nrow(Science)
unweighted <- mirt(Science, 1)
weighted <- mirt(Science, 1, survey.weights=survey.weights)

#####
# empirical dimensionality testing that includes 'guessing'

data(SAT12)
data <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
itemstats(data)

mod1 <- mirt(data, 1)
extract.mirt(mod1, 'time') #time elapsed for each estimation component

# optionally use Newton-Raphson for (generally) faster convergence in the M-step's
mod1 <- mirt(data, 1, optimizer = 'NR')
extract.mirt(mod1, 'time')

mod2 <- mirt(data, 2, optimizer = 'NR')
# difficulty converging with reduced quadpts, reduce TOL
mod3 <- mirt(data, 3, TOL = .001, optimizer = 'NR')
anova(mod1,mod2)
anova(mod2, mod3) #negative AIC, 2 factors probably best

# same as above, but using the QMCEM method for generally better accuracy in mod3

```

```

mod3 <- mirt(data, 3, method = 'QMCEM', TOL = .001, optimizer = 'NR')
anova(mod2, mod3)

# with fixed guessing parameters
mod1g <- mirt(data, 1, guess = .1)
coef(mod1g)

#####
# graded rating scale example

# make some data
set.seed(1234)
a <- matrix(rep(1, 10))
d <- matrix(c(1,0.5,-.5,-1), 10, 4, byrow = TRUE)
c <- seq(-1, 1, length.out=10)
data <- simdata(a, d + c, 2000, itemtype = rep('graded',10))
itemstats(data)

mod1 <- mirt(data, 1)
mod2 <- mirt(data, 1, itemtype = 'grsm')
coef(mod2)
anova(mod2, mod1) #not sig, mod2 should be preferred
itemplot(mod2, 1)
itemplot(mod2, 5)
itemplot(mod2, 10)

#####
# 2PL nominal response model example (Suh and Bolt, 2010)
data(SAT12)
SAT12[SAT12 == 8] <- NA #set 8 as a missing value
head(SAT12)

# correct answer key
key <- c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5)
scoredSAT12 <- key2binary(SAT12, key)
mod0 <- mirt(scoredSAT12, 1)

# for first 5 items use 2PLNRM and nominal
scoredSAT12[,1:5] <- as.matrix(SAT12[,1:5])
mod1 <- mirt(scoredSAT12, 1, c(rep('nominal',5),rep('2PL', 27)))
mod2 <- mirt(scoredSAT12, 1, c(rep('2PLNRM',5),rep('2PL', 27)), key=key)
coef(mod0)$Item.1
coef(mod1)$Item.1
coef(mod2)$Item.1
itemplot(mod0, 1)
itemplot(mod1, 1)
itemplot(mod2, 1)

# compare added information from distractors
Theta <- matrix(seq(-4,4,.01))
par(mfrow = c(2,3))
for(i in 1:5){
  info <- iteminfo(extract.item(mod0,i), Theta)

```

```

    info2 <- iteminfo(extract.item(mod2,i), Theta)
    plot(Theta, info2, type = 'l', main = paste('Information for item', i), ylab = 'Information')
    lines(Theta, info, col = 'red')
  }
  par(mfrow = c(1,1))

# test information
plot(Theta, testinfo(mod2, Theta), type = 'l', main = 'Test information', ylab = 'Information')
lines(Theta, testinfo(mod0, Theta), col = 'red')

#####
# using the MH-RM algorithm
data(LSAT7)
fulldata <- expand.table(LSAT7)
(mod1 <- mirt(fulldata, 1, method = 'MHRM'))

# Confirmatory models

# simulate data
a <- matrix(c(
  1.5,NA,
  0.5,NA,
  1.0,NA,
  1.0,0.5,
  NA,1.5,
  NA,0.5,
  NA,1.0,
  NA,1.0),ncol=2,byrow=TRUE)

d <- matrix(c(
  -1.0,NA,NA,
  -1.5,NA,NA,
  1.5,NA,NA,
  0.0,NA,NA,
  3.0,2.0,-0.5,
  2.5,1.0,-1,
  2.0,0.0,NA,
  1.0,NA,NA),ncol=3,byrow=TRUE)

sigma <- diag(2)
sigma[1,2] <- sigma[2,1] <- .4
items <- c(rep('2PL',4), rep('graded',3), '2PL')
dataset <- simdata(a,d,2000,items,sigma)

# analyses
# CIFA for 2 factor crossed structure

model.1 <- '
  F1 = 1-4
  F2 = 4-8
  COV = F1*F2'

# compute model, and use parallel computation of the log-likelihood

```

```

if(interactive()) mirtCluster()
mod1 <- mirt(dataset, model.1, method = 'MHRM')
coef(mod1)
summary(mod1)
residuals(mod1)

#####
# bifactor
model.3 <- '
  G = 1-8
  F1 = 1-4
  F2 = 5-8'

mod3 <- mirt(dataset,model.3, method = 'MHRM')
coef(mod3)
summary(mod3)
residuals(mod3)
anova(mod1,mod3)

#####
# polynomial/combinations
data(SAT12)
data <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))

model.quad <- '
  F1 = 1-32
  (F1*F1) = 1-32'

model.combo <- '
  F1 = 1-16
  F2 = 17-32
  (F1*F2) = 1-8'

(mod.quad <- mirt(data, model.quad))
summary(mod.quad)
(mod.combo <- mirt(data, model.combo))
anova(mod.combo, mod.quad)

# non-linear item and test plots
plot(mod.quad)
plot(mod.combo, type = 'SE')
itemplot(mod.quad, 1, type = 'score')
itemplot(mod.combo, 2, type = 'score')
itemplot(mod.combo, 2, type = 'infocontour')

## empirical histogram examples (normal, skew and bimodality)
# make some data
set.seed(1234)
a <- matrix(rlnorm(50, .2, .2))
d <- matrix(rnorm(50))
ThetaNormal <- matrix(rnorm(2000))

```

```

ThetaBimodal <- scale(matrix(c(rnorm(1000, -2), rnorm(1000,2)))) #bimodal
ThetaSkew <- scale(matrix(rchisq(2000, 3))) #positive skew
datNormal <- simdata(a, d, 2000, itemtype = '2PL', Theta=ThetaNormal)
datBimodal <- simdata(a, d, 2000, itemtype = '2PL', Theta=ThetaBimodal)
datSkew <- simdata(a, d, 2000, itemtype = '2PL', Theta=ThetaSkew)

normal <- mirt(datNormal, 1, dentype = "empiricalhist")
plot(normal, type = 'empiricalhist')
histogram(ThetaNormal, breaks=30)

bimodal <- mirt(datBimodal, 1, dentype = "empiricalhist")
plot(bimodal, type = 'empiricalhist')
histogram(ThetaBimodal, breaks=30)

skew <- mirt(datSkew, 1, dentype = "empiricalhist")
plot(skew, type = 'empiricalhist')
histogram(ThetaSkew, breaks=30)

#####
# non-linear parameter constraints with Rsolnp package (nloptr supported as well):
# Find Rasch model subject to the constraint that the intercepts sum to 0

dat <- expand.table(LSAT6)
itemstats(dat)

# free latent mean and variance terms
model <- 'Theta = 1-5
        MEAN = Theta
        COV = Theta*Theta'

# view how vector of parameters is organized internally
sv <- mirt(dat, model, itemtype = 'Rasch', pars = 'values')
sv[sv$est, ]

# constraint: create function for solnp to compute constraint, and declare value in eqB
eqfun <- function(p, optim_args) sum(p[1:5]) #could use browser() here, if it helps
LB <- c(rep(-15, 6), 1e-4) # more reasonable lower bound for variance term

mod <- mirt(dat, model, sv=sv, itemtype = 'Rasch', optimizer = 'solnp',
  solnp_args=list(eqfun=eqfun, eqB=0, LB=LB))
print(mod)
coef(mod)
(ds <- sapply(coef(mod)[1:5], function(x) x[, 'd']))
sum(ds)

# same likelihood location as: mirt(dat, 1, itemtype = 'Rasch')

#####
# latent regression Rasch model

# simulate data
set.seed(1234)

```

```

N <- 1000

# covariates
X1 <- rnorm(N); X2 <- rnorm(N)
covdata <- data.frame(X1, X2, X3 = rnorm(N))
Theta <- matrix(0.5 * X1 + -1 * X2 + rnorm(N, sd = 0.5))

# items and response data
a <- matrix(1, 20); d <- matrix(rnorm(20))
dat <- simdata(a, d, 1000, itemtype = '2PL', Theta=Theta)

# unconditional Rasch model
mod0 <- mirt(dat, 1, 'Rasch', SE=TRUE)
coef(mod0, printSE=TRUE)

# conditional model using X1, X2, and X3 (bad) as predictors of Theta
mod1 <- mirt(dat, 1, 'Rasch', covdata=covdata, formula = ~ X1 + X2 + X3, SE=TRUE)
coef(mod1, printSE=TRUE)
coef(mod1, simplify=TRUE)
anova(mod0, mod1) # jointly significant predictors of theta

# large sample z-ratios and p-values (if one cares)
cfs <- coef(mod1, printSE=TRUE)
(z <- cfs$lr.betas[[1]] / cfs$lr.betas[[2]])
round(pnorm(abs(z[,1]), lower.tail=FALSE)*2, 3)

# drop predictor for nested comparison
mod1b <- mirt(dat, 1, 'Rasch', covdata=covdata, formula = ~ X1 + X2)
anova(mod1b, mod1)

# compare to mixedmirt() version of the same model
mod1.mixed <- mixedmirt(dat, 1, itemtype='Rasch',
                        covdata=covdata, lr.fixed = ~ X1 + X2 + X3, SE=TRUE)
coef(mod1.mixed)
coef(mod1.mixed, printSE=TRUE)

# draw plausible values for secondary analyses
pv <- fscores(mod1, plausible.draws = 10)
pvmods <- lapply(pv, function(x, covdata) lm(x ~ covdata$X1 + covdata$X2),
                 covdata=covdata)
# population characteristics recovered well, and can be averaged over
so <- lapply(pvmods, summary)
so

# compute Rubin's multiple imputation average
par <- lapply(so, function(x) x$coefficients[, 'Estimate'])
SEpar <- lapply(so, function(x) x$coefficients[, 'Std. Error'])
averageMI(par, SEpar)

#####
# Example using Gauss-Hermite quadrature with custom input functions

library(fastGHQuad)

```

```

data(SAT12)
data <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
GH <- gaussHermiteData(50)
Theta <- matrix(GH$x)

# This prior works for uni- and multi-dimensional models
prior <- function(Theta, Etable){
  P <- grid <- GH$w / sqrt(pi)
  if(ncol(Theta) > 1)
    for(i in 2:ncol(Theta))
      P <- expand.grid(P, grid)
  if(!is.vector(P)) P <- apply(P, 1, prod)
  P
}

GHmod1 <- mirt(data, 1, optimizer = 'NR',
  technical = list(customTheta = Theta, customPriorFun = prior))
coef(GHmod1, simplify=TRUE)

Theta2 <- as.matrix(expand.grid(Theta, Theta))
GHmod2 <- mirt(data, 2, optimizer = 'NR', TOL = .0002,
  technical = list(customTheta = Theta2, customPriorFun = prior))
summary(GHmod2, suppress=.2)

#####
# Davidian curve example

dat <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
dav <- mirt(dat, 1, dentype = 'Davidian-4') # use four smoothing parameters
plot(dav, type = 'Davidian') # shape of latent trait distribution
coef(dav, simplify=TRUE)

fs <- fscores(dav) # assume normal prior
fs2 <- fscores(dav, use_dentype_estimate=TRUE) # use Davidian estimated prior shape
head(cbind(fs, fs2))

itemfit(dav) # assume normal prior
itemfit(dav, use_dentype_estimate=TRUE) # use Davidian estimated prior shape

#####
# Unfolding models

# polytomous hyperbolic cosine model with
# estimated latitude of acceptance (rho parameters)
mod <- mirt(Science, model=1, itemtype = 'hcm')
coef(mod, simplify=TRUE)$items
coef(mod, simplify=TRUE, IRTpars=TRUE)$items

plot(mod)
plot(mod, type = 'trace')
plot(mod, type = 'itemscore')

```

```

# EAP estimates
fs <- fscores(mod)
head(fs)

itemfit(mod)
M2(mod, type = 'C2')

#####
# 5PL and restricted 5PL example
dat <- expand.table(LSAT7)

mod2PL <- mirt(dat)
mod2PL

# Following does not converge without including strong priors
# mod5PL <- mirt(dat, itemtype = '5PL')
# mod5PL

# restricted version of 5PL (asymmetric 2PL)
model <- 'Theta = 1-5
        FIXED = (1-5, g), (1-5, u)'

mod2PL_asym <- mirt(dat, model=model, itemtype = '5PL')
mod2PL_asym
coef(mod2PL_asym, simplify=TRUE)
coef(mod2PL_asym, simplify=TRUE, IRTpars=TRUE)

# no big difference statistically or visually
anova(mod2PL, mod2PL_asym)
plot(mod2PL, type = 'trace')
plot(mod2PL_asym, type = 'trace')

#####
# LLTM example

a <- matrix(rep(1,30))
d <- rep(c(1,0, -1),each = 10) # first easy, then medium, last difficult
dat <- simdata(a, d, 1000, itemtype = '2PL')

# unconditional model for intercept comparisons
mod <- mirt(dat, itemtype = 'Rasch')
coef(mod, simplify=TRUE)

# Suppose that the first 10 items were suspected to be easy, followed by 10 medium difficulty items,
# then finally the last 10 items are difficult,
# and we wish to test this item structure hypothesis (more intercept designs are possible
# by including more columns).
itemdesign <- data.frame(difficulty =
  factor(c(rep('easy', 10), rep('medium', 10), rep('hard', 10))))
rownames(itemdesign) <- colnames(dat)
itemdesign

```

```

# LLTM with mirt()
lltm <- mirt(dat, itemtype = 'Rasch', SE=TRUE,
  item.formula = ~ 0 + difficulty, itemdesign=itemdesign)
coef(lltm, simplify=TRUE)
coef(lltm, printSE=TRUE)
anova(lltm, mod) # models fit effectively the same; hence, intercept variability well captured

# additional information for LLTM
plot(lltm)
plot(lltm, type = 'trace')
itemplot(lltm, item=1)
itemfit(lltm)
head(fscores(lltm)) #EAP estimates
fscores(lltm, method='EAPsum', full.scores=FALSE)
M2(lltm) # goodness of fit
head(personfit(lltm))
residuals(lltm)

# intercept across items also possible by removing ~ 0 portion, just interpreted differently
lltm.int <- mirt(dat, itemtype = 'Rasch',
  item.formula = ~ difficulty, itemdesign=itemdesign)
anova(lltm, lltm.int) # same
coef(lltm.int, simplify=TRUE)

# using unconditional modeling for first four items
itemdesign.sub <- itemdesign[5:nrow(itemdesign), , drop=FALSE]
itemdesign.sub # note that rownames are required in this case
lltm.4 <- mirt(dat, itemtype = 'Rasch',
  item.formula = ~ 0 + difficulty, itemdesign=itemdesign.sub)
coef(lltm.4, simplify=TRUE) # first four items are the standard Rasch
anova(lltm, lltm.4) # similar fit, hence more constrained model preferred

# LLTM with mixedmirt() (more flexible in general, but slower)
LLTM <- mixedmirt(dat, model=1, fixed = ~ 0 + difficulty,
  itemdesign=itemdesign, SE=FALSE)
summary(LLTM)
coef(LLTM)

# LLTM with random error estimate (not supported with mirt() )
LLTM.e <- mixedmirt(dat, model=1, fixed = ~ 0 + difficulty,
  random = ~ 1|items, itemdesign=itemdesign, SE=FALSE)
coef(LLTM.e)

#####
# General MLTM example (Embretson, 1984)

set.seed(42)

as <- matrix(rep(1,60), ncol=2)
as[11:18,1] <- as[1:9,2] <- 0
d1 <- rep(c(3,1),each = 6) # first easy, then medium, last difficult for first trait

```

```

d2 <- rep(c(0,1,2),times = 4)    # difficult to easy
d <- rnorm(18)
ds <- rbind(cbind(d1=NA, d2=d), cbind(d1, d2))
(pars <- data.frame(a=as, d=ds))
dat <- simdata(as, ds, 2500,
  itemtype = c(rep('dich', 18), rep('partcomp', 12)))
itemstats(dat)

# unconditional model
syntax <- "theta1 = 1-9, 19-30
          theta2 = 10-30
          COV = theta1*theta2"
itemtype <- c(rep('Rasch', 18), rep('PC1PL', 12))
mod <- mirt(dat, syntax, itemtype=itemtype)
coef(mod, simplify=TRUE)
data.frame(est=coef(mod, simplify=TRUE)$items, pop=data.frame(a=as, d=ds))
itemplot(mod, 1)
itemplot(mod, 30)

# MLTM design only for PC1PL items
itemdesign <- data.frame(t1_difficulty= factor(d1, labels=c('medium', 'easy')),
  t2_difficulty=factor(d2, labels=c('hard', 'medium', 'easy')))
rownames(itemdesign) <- colnames(dat)[19:30]
itemdesign

# fit MLTM design, leaving first 18 items as 'Rasch' type
mltm <- mirt(dat, syntax, itemtype=itemtype, itemdesign=itemdesign,
  item.formula = list(theta1 ~ 0 + t1_difficulty,
    theta2 ~ 0 + t2_difficulty), SE=TRUE)

coef(mltm, simplify=TRUE)
coef(mltm, printSE=TRUE)
anova(mltm, mod) # similar fit; hence more constrained version preferred
M2(mltm) # goodness of fit
head(personfit(mltm))
residuals(mltm)

# EAP estimates
fscores(mltm) |> head()

## End(Not run)

```

mirt.model

Specify model information

Description

The `mirt.model` function scans/reads user input to specify the confirmatory model. Item locations must be used in the specifications if no `itemnames` argument is supplied. This is called implicitly by estimation functions when a string is passed to the `model` argument.

Usage

```
mirt.model(
  input = NULL,
  itemnames = NULL,
  file = "",
  COV = NULL,
  quiet = TRUE,
  ...
)
```

Arguments

<code>input</code>	input for writing out the model syntax. Can either be a string declaration of class character or the so-called Q-matrix or class <code>matrix</code> that specifies the model either with integer or logical values. If the Q-matrix method is chosen covariances terms can be specified with the COV input
<code>itemnames</code>	a character vector or factor indicating the item names. If a data.frame or matrix object is supplied the names will be extracted using <code>colnames(itemnames)</code> . Supplying this input allows the syntax to be specified with the raw item names rather than item locations
<code>file</code>	a input specifying an external file that declares the input.
<code>COV</code>	a symmetric, logical matrix used to declare which covariance terms are estimated
<code>quiet</code>	logical argument passed to <code>scan()</code> to suppress console read message
<code>...</code>	additional arguments for <code>scan()</code>

Details

Factors are first named and then specify which numerical items they affect (i.e., where the slope is not equal to 0), separated either by commas or by - to indicate a range of items. Products between factors may be specified by enclosing the left hand term within brackets. To finish the declaration of a model simply enter a blank line with only a carriage return (i.e., the 'enter' or 'return' key), or instead read in an input version of the model syntax. The associated slopes throughout the package label these coefficients as $a_1, a_2, \dots, a_k$, where the associated number is assigned according to the respective order of the defined factors.

For example, if the syntax were

```
"G = 1-10 F = 1-5 A = 6-10"
```

then the G factor would be assigned the slopes a_1 for each item, F assigned the slopes a_2 , and A assigned the slopes a_3 . The same principle applies to the `bfactor` function whereby the slopes are automatically included for the specific factors after the general factor structure has been assigned.

There is an optional keyword for specifying the correlation between relationships between factors called COV, and non-linear factor products can be included by enclosing the product combination on the left hand side of the declaration (e.g., $(F1 \times F1)$ would create a quadratic factor for F1).

The keywords `CONSTRAIN`, `CONSTRAINB`, `PRIOR`, `FIXED`, `FREE`, `START`, `UBOUND`, `LBOUND` can be applied to specific sub-groups in multiple-group models by included square brackets before the =

sign, where groups are separated by commas. For example, to apply within-group equality constraints to a group called "male", then specifying:

```
CONSTRAIN [male] = (1-5, a1)
```

is appropriate, while specifying the same constraints to the sub-groups "male" and "female" would appear as

```
CONSTRAIN [male, female] = (1-5, a1)
```

For all other groups in the multi-group model, these within-group equality constraints would not appear. Therefore, these bracketed group specifications are useful when modifying priors, starting values, between/within group equality constraints, and so on when the specifications for each sub-group may differ.

Additionally, the use of negations can be used to omit specific groups in the constraint specifications by prefixing the string with a - operator, such as the following which applies between-group constraints to all groups except "Group2" and "Group3":

```
CONSTRAINB [-Group2, -Group3] = (1-5, a1)
```

Finally, the keyword GROUP can be used to specify the group-level hyper-parameter terms, such as the means and variance of the default Gaussian distribution. For example, to set the starting value of the variance parameter (COV_11) to 1.5:

```
START = (GROUP, COV_11, 1.5)
```

COV Specify the relationship between the latent factors. Estimating a correlation between factors is declared by joining the two factors with an asterisk (e.g., F1*F2), or with an asterisk between three or more factors to estimate all the possible correlations (e.g., F1*F2*F3). Specifications with the same factor (e.g., F1*F1) will free the variance of said factor instead

MEAN A comma separated list specifying which latent factor means to freely estimate. E.g., MEAN = F1, F2 will free the latent means for factors F1 and F2

CONSTRAIN A bracketed, comma separated list specifying equality constraints between items. The input format is CONSTRAIN = (items, ..., parameterName(s)), (items, ..., parameterName).

For example, in a single group 10-item dichotomous tests, using the default 2PL model, the first and last 5 item slopes (a1) can be constrained to be equal by using CONSTRAIN = (1-5, a1), (6-10, a1), or some combination such as CONSTRAIN = (1-3, 4, 5, a1), (6, 7, 8-10, a1).

When constraining parameters to be equal across items with different parameter names, a balanced bracketed vector must be supplied. E.g., setting the first slope for item 1 equal to the second slope in item 3 would be CONSTRAIN = (1, 3, a1, a2)

CONSTRAINB A bracketed, comma separate list specifying equality constraints between groups. The input format is CONSTRAINB = (items, ..., parameterName), (items, ..., parameterName).

For example, in a two group 10-item dichotomous tests, using the default 2PL model, the first 5 item slopes (a1) can be constrained to be equal across both groups by using CONSTRAINB = (1-5, a1), or some combination such as CONSTRAINB = (1-3, 4, 5, a1)

PRIOR A bracketed, comma separate list specifying prior parameter distributions. The input format is PRIOR = (items, ..., parameterName, priorType, val1, val2), (items, ..., parameterName, priorType, val1, val2). For example, in a single group 10-item dichotomous tests, using the default 2PL model, defining a normal prior of N(0,2) for the first 5 item intercepts (d) can be defined by PRIOR = (1-5, d, norm, 0, 2)

Currently supported priors are of the form: (items, norm, mean, sd) for the normal/Gaussian, (items, lnorm, log_mean, log_sd) for log-normal, (items, beta, alpha, beta) for beta,

and (items, expbeta, alpha, beta) for the beta distribution after applying the function `plogis` to the input value (note, this is specifically for applying a beta prior to the lower-bound parameters in 3/4PL models). Note that the beta prior specification focuses on the mode of the beta distribution centered around the value $\frac{\alpha}{(\alpha+\beta)}$ rather than the mean; for the mean of the beta distribution specify alpha + 1 and beta + 1 instead

LBOUND A bracketed, comma separate list specifying lower bounds for estimated parameters (used in optimizers such as L-BFGS-B and nlminb). The input format is LBOUND = (items, ..., parameterName, value), (items, ..., parameterName, value).

For example, in a single group 10-item dichotomous tests, using the 3PL model and setting lower bounds for the 'g' parameters for the first 5 items to 0.2 is accomplished with LBOUND = (1-5, g, 0.2)

UBOUND same as LBOUND, but specifying upper bounds in estimated parameters

START A bracketed, comma separate list specifying the starting values for individual parameters. The input is of the form (items, ..., parameterName, value). For instance, setting the 10th and 12th to 15th item slope parameters (a1) to 1.0 is specified with START = (10, 12-15, a1, 1.0)

For more hands on control of the starting values pass the argument pars = 'values' through whatever estimation function is being used

FIXED A bracketed, comma separate list specifying which parameters should be fixed at their starting values (i.e., not freely estimated). The input is of the form (items, ..., parameterName). For instance, fixing the 10th and 12th to 15th item slope parameters (a1) is accomplished with FIXED = (10, 12-15, a1)

For more hands on control of the estimated values pass the argument pars = 'values' through whatever estimation function is being used

FREE Equivalent to the FIXED input, except that parameters are freely estimated instead of fixed at their starting value

NEXPLORE Number of exploratory factors to extract. Usually this is not required because passing a numeric value to the model argument in the estimation function will generate an exploratory factor analysis model, however if different start values, priors, lower and upper bounds, etc, are desired then this input can be used

Value

Returns a model specification object to be used in `mirt`, `bfactor`, `multipleGroup`, or `mixedmirt`

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com> and Alexander Robitzsch

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
## Not run:

# interactively through the console (not run)
#model <- mirt.model()
# F1 = 1,2,3,4-10
# F2 = 10-20
# (F1*F2) = 1,2,3,4-10
# COV = F1*F2

# Or alternatively with a string input
s <- 'F1 = 1,2,3,4-10
      F2 = 10-20
      (F1*F2) = 1,2,3,4-10
      COV = F1*F2'
model <- mirt.model(s)

# strings can also be passed to the estimation functions directly,
# which silently calls mirt.model(). E.g., using the string above:
# mod <- mirt(data, s)

# Q-matrix specification
Q <- matrix(c(1,1,1,0,0,0,0,0,0,1,1,1), ncol=2, dimnames = list(NULL, c('Factor1', 'Factor2')))
COV <- matrix(c(FALSE, TRUE, TRUE, FALSE), 2)
model <- mirt.model(Q, COV=COV)

## constrain various items slopes and all intercepts in single group model to be equal,
# and use a log-normal prior for all the slopes
s <- 'F = 1-10
      CONSTRAIN = (1-3, 5, 6, a1), (1-10, d)
      PRIOR = (1-10, a1, lnorm, .2, .2)'
model <- mirt.model(s)

## constrain various items slopes and intercepts across groups for use in multipleGroup(),
# and constrain first two slopes within 'group1' to be equal
s <- 'F = 1-10
      CONSTRAIN = (1-2, a1)
      CONSTRAINB = (1-3, 5, 6, a1), (1-10, d)'
model <- mirt.model(s)

## specify model using raw item names
data(data.read, package = 'sirt')
dat <- data.read

# syntax with variable names
mirtsyn2 <- "
      F1 = A1,B2,B3,C4
      F2 = A1-A4,C2,C4
```

```

      MEAN = F1
      COV = F1*F1, F1*F2
      CONSTRAIN=(A2-A4,a2),(A3,C2,d)
      PRIOR = (C3,A2-A4,a2,lnorm, .2, .2),(B3,d,norm,0,.0001)"
# create a mirt model
mirtmodel <- mirt.model(mirtsyn2, itemnames=dat)
# or equivalently:
# mirtmodel <- mirt.model(mirtsyn2, itemnames=colnames(dat))

# mod <- mirt(dat , mirtmodel)

# using sprintf() to functionally fill in information (useful for long tests
# or more complex specifications)
nitems <- 100
s <- sprintf('F = 1-%i
      CONSTRAIN = (%s, a1)
      CONSTRAINB = (%s, a1), (1-%i, d)',
      nitems, "1,2,4,50,100",
      paste0(1:45, collapse=','),
      nitems)
cat(s)
model <- mirt.model(s)

## End(Not run)

```

mirtCluster

Define a parallel cluster object to be used in internal functions

Description

This function defines a object that is placed in a relevant internal environment defined in mirt. Internal functions such as calcLogLik, fscores, etc, will utilize this object automatically to capitalize on parallel processing architecture. The object defined is a call from `parallel::makeCluster()`. Note that if you are defining other parallel objects (for simulation designs, for example) it is not recommended to define a mirtCluster.

Usage

```
mirtCluster(spec, omp_threads, remove = FALSE, ...)
```

Arguments

spec	input that is passed to <code>parallel::makeCluster()</code> . If no input is given the maximum number of available local cores minus 1 will be used. Setting this to NULL will skip a new definition (allows <code>omp_threads</code> to be used independently)
omp_threads	number of OpenMP threads to use (currently applies to E-step computations only). Not used when argument input is missing

remove logical; remove previously defined mirtCluster()?
 ... additional arguments to pass to parallel::makeCluster

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```
## Not run:
if(interactive()){
  # use all available cores
  mirtCluster()
  mirtCluster(remove = TRUE)

  # make 4 cores available for parallel computing
  mirtCluster(4)
  mirtCluster(remove = TRUE)

  # create 3 core architecture in R, and 4 thread architecture with OpenMP
  mirtCluster(spec = 3, omp_threads = 4)

  # leave previous multicore objects, but change omp_threads
  mirtCluster(spec = NULL, omp_threads = 2)
}

## End(Not run)
```

MixedClass-class	Class "MixedClass"
------------------	--------------------

Description

Defines the object returned from [mixedmirt](#).

Slots

Call: function call
 Data: list of data, sometimes in different forms
 Options: list of estimation options
 Fit: a list of fit information
 Model: a list of model-based information

ParObjects: a list of the S4 objects used during estimation

OptimInfo: a list of arguments from the optimization process

Internals: a list of internal arguments for secondary computations (inspecting this object is generally not required)

vcov: a matrix represented the asymptotic covariance matrix of the parameter estimates

time: a data.frame indicating the breakdown of computation times in seconds

Methods

coef signature(object = "MixedClass")

print signature(x = "MixedClass")

residuals signature(object = "MixedClass")

show signature(object = "MixedClass")

summary signature(object = "MixedClass")

logLik signature(object = "MixedClass")

anova signature(object = "MixedClass")

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

mixedmirt

Mixed effects modeling for MIRT models

Description

mixedmirt fits MIRT models using FIML estimation to dichotomous and polytomous IRT models conditional on fixed and random effect of person and item level covariates. This can also be understood as 'explanatory IRT' if only fixed effects are modeled, or multilevel/mixed IRT if random and fixed effects are included. The method uses the MH-RM algorithm exclusively. Additionally, computation of the log-likelihood can be sped up by using parallel estimation via [mirtCluster](#).

Usage

```

mixedmirt(
  data,
  covdata = NULL,
  model = 1,
  fixed = ~1,
  random = NULL,
  itemtype = "Rasch",
  lr.fixed = ~1,
  lr.random = NULL,
  itemdesign = NULL,
  constrain = NULL,
  pars = NULL,
  return.design = FALSE,
  SE = TRUE,
  internal_constraints = TRUE,
  technical = list(SETol = 1e-04),
  ...
)

```

Arguments

<code>data</code>	a matrix or data.frame that consists of numerically ordered data, organized in the form of integers, with missing data coded as NA
<code>covdata</code>	a data.frame that consists of the <code>nrow(data)</code> by <code>K</code> 'person level' fixed and random predictors. If missing data are present in this object then the observations from <code>covdata</code> and <code>data</code> will be removed row-wise via the <code>rowSums(is.na(covdata)) > 0</code>
<code>model</code>	an object returned from, or a string to be passed to, <code>mirt.model()</code> to declare how the IRT model is to be estimated. See mirt.model and mirt for more detail
<code>fixed</code>	a right sided R formula for specifying the fixed effect (aka 'explanatory') predictors from <code>covdata</code> and <code>itemdesign</code> . To estimate the intercepts for each item the keyword <code>items</code> is reserved and automatically added to the <code>itemdesign</code> input. If any polytomous items are being model the <code>items</code> argument is not valid since all intercept parameters are freely estimated and identified with the parameterizations found in mirt , and the first column in the fixed design matrix (commonly the intercept or a reference group) is omitted
<code>random</code>	a right sided formula or list of formulas containing crossed random effects of the form <code>v1 + ... v_n G</code> , where <code>G</code> is the grouping variable and <code>v_n</code> are random numeric predictors within each group. If no intercept value is specified then by default the correlations between the <code>v</code> 's and <code>G</code> are estimated, but can be suppressed by including the <code>~-1 + ...</code> or <code>0</code> constant. <code>G</code> may contain interaction terms, such as <code>group:items</code> to include cross or person-level interactions effects
<code>itemtype</code>	same as <code>itemtype</code> in mirt , except when the fixed or random inputs are used does not support the following item types: <code>c('PC2PL', 'PC3PL', '2PLNRM', '3PLNRM', '3PLuNRM', '4PLNRM')</code>

<code>lr.fixed</code>	an R formula (or list of formulas) to specify regression effects in the latent variables from the variables in <code>covdata</code> . This is used to construct models such as the so-called 'latent regression model' to explain person-level ability/trait differences. If a named list of formulas is supplied (where the names correspond to the latent trait names in <code>model</code>) then specific regression effects can be estimated for each factor. Supplying a single formula will estimate the regression parameters for all latent traits by default.
<code>lr.random</code>	a list of random effect terms for modeling variability in the latent trait scores, where the syntax uses the same style as in the <code>random</code> argument. Useful for building so-called 'multilevel IRT' models which are non-Rasch (multilevel Rasch models do not technically require these because they can be built using the fixed and random inputs alone)
<code>itemdesign</code>	a <code>data.frame</code> object used to create a design matrix for the items, where each <code>nrow(itemdesign) == nitems</code> and the number of columns is equal to the number of fixed effect predictors (i.e., item intercepts). By default an <code>items</code> variable is reserved for modeling the item intercept parameters
<code>constrain</code>	a list indicating parameter equality constraints. See mirt for more detail
<code>pars</code>	used for parameter starting values. See mirt for more detail
<code>return.design</code>	logical; return the design matrices before they have (potentially) been reassigned?
<code>SE</code>	logical; compute the standard errors by approximating the information matrix using the MHRM algorithm? Default is TRUE
<code>internal_constraints</code>	logical; use the internally defined constraints for constraining effects across persons and items? Default is TRUE. Setting this to FALSE runs the risk of under-identification
<code>technical</code>	the technical list passed to the MH-RM estimation engine, with the <code>SEtol</code> default increased to .0001. Additionally, the argument <code>RANDSTART</code> is available to indicate at which iteration (during the burn-in stage) the additional random effect variables should begin to be approximated (i.e., elements in <code>lr.random</code> and <code>random</code>). The default for <code>RANDSTART</code> is to start at iteration 100, and when random effects are included the default number of burn-in iterations is increased from 150 to 200. See mirt for further details
<code>...</code>	additional arguments to be passed to the MH-RM estimation engine. See mirt for more details and examples

Details

For dichotomous response models, `mixedmirt` follows the general form

$$P(x = 1|\theta, \psi) = g + \frac{(u - g)}{1 + \exp(-1 * [\theta a + X\beta + Z\delta])}$$

where X is a design matrix with associated β fixed effect intercept coefficients, and Z is a design matrix with associated δ random effects for the intercepts. For simplicity and easier interpretation, the unique item intercept values typically found in $X\beta$ are extracted and reassigned within

mirt's 'intercept' parameters (e.g., 'd'). To observe how the design matrices are structured prior to reassignment and estimation pass the argument `return.design = TRUE`.

Polytomous IRT models follow a similar format except the item intercepts are automatically estimated internally, rendering the `items` argument in the fixed formula redundant and therefore must be omitted from the specification. If there are a mixture of dichotomous and polytomous items the intercepts for the dichotomous models are also estimated for consistency.

The decomposition of the θ parameters is also possible to form latent regression and multilevel IRT models by using the `lr.fixed` and `lr.random` inputs. These effects decompose θ such that

$$\theta = V\Gamma + W\zeta + \epsilon$$

where V and W are fixed and random effects design matrices for the associated coefficients.

To simulate expected a posteriori predictions for the random effect terms use the [randef](#) function.

Value

function returns an object of class `MixedClass` ([MixedClass-class](#)).

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Chalmers, R. P. (2015). Extended Mixed-Effects Item Response Models with the MH-RM Algorithm. *Journal of Educational Measurement*, 52, 200-222. doi:10.1111/jedm.12072

See Also

[mirt](#), [randef](#), [fixef](#), [boot.mirt](#)

Examples

```
## Not run:

# make some data
set.seed(1234)
N <- 750
a <- matrix(rlnorm(10,.3,1),10,1)
d <- matrix(rnorm(10), 10)
Theta <- matrix(sort(rnorm(N)))
pseudoIQ <- Theta * 5 + 100 + rnorm(N, 0, 5)
pseudoIQ <- (pseudoIQ - mean(pseudoIQ))/10 #rescale variable for numerical stability
group <- factor(rep(c('G1','G2','G3'), each = N/3))
data <- simdata(a,d,N, itemtype = rep('2PL',10), Theta=Theta)
covdata <- data.frame(group, pseudoIQ)
```

```

itemstats(data)

# use parallel computing
if(interactive()) mirtCluster()

# specify IRT model
model <- 'Theta = 1-10'

# model with no person predictors
mod0 <- mirt(data, model, itemtype = 'Rasch')

# group as a fixed effect predictor (aka, uniform dif)
mod1 <- mixedmirt(data, covdata, model, fixed = ~ 0 + group + items)
anova(mod0, mod1)
summary(mod1)
coef(mod1)

# same model as above in lme4
wide <- data.frame(id=1:nrow(data),data,covdata)
long <- reshape2::melt(wide, id.vars = c('id', 'group', 'pseudoIQ'))
library(lme4)
lmod0 <- glmer(value ~ 0 + variable + (1|id), long, family = binomial)
lmod1 <- glmer(value ~ 0 + group + variable + (1|id), long, family = binomial)
anova(lmod0, lmod1)

# model using 2PL items instead of Rasch
mod1b <- mixedmirt(data, covdata, model, fixed = ~ 0 + group + items, itemtype = '2PL')
anova(mod1, mod1b) #better with 2PL models using all criteria (as expected, given simdata pars)

# continuous predictor with group
mod2 <- mixedmirt(data, covdata, model, fixed = ~ 0 + group + items + pseudoIQ)
summary(mod2)
anova(mod1b, mod2)

# view fixed design matrix with and without unique item level intercepts
withinint <- mixedmirt(data, covdata, model, fixed = ~ 0 + items + group, return.design = TRUE)
withoutint <- mixedmirt(data, covdata, model, fixed = ~ 0 + group, return.design = TRUE)

# notice that in result above, the intercepts 'items1 to items 10' were reassigned to 'd'
head(withinint$X)
tail(withinint$X)
head(withoutint$X) # no intercepts design here to be reassigned into item intercepts
tail(withoutint$X)

#####
### random effects
# make the number of groups much larger
covdata$group <- factor(rep(paste0('G',1:50), each = N/50))

# random groups
rmod1 <- mixedmirt(data, covdata, 1, fixed = ~ 0 + items, random = ~ 1|group)
summary(rmod1)
coef(rmod1)

```

```

# random groups and random items
rmod2 <- mixedmirt(data, covdata, 1, random = list(~ 1|group, ~ 1|items))
summary(rmod2)
eff <- randef(rmod2) #estimate random effects

# random slopes with fixed intercepts (suppressed correlation)
rmod3 <- mixedmirt(data, covdata, 1, fixed = ~ 0 + items, random = ~ -1 + pseudoIQ|group)
summary(rmod3)
eff <- randef(rmod3)
str(eff)

#####
## LLTM, and 2PL version of LLTM
data(SAT12)
data <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
model <- 'Theta = 1-32'

# for unconditional intercept comparison
mod <- mirt(data, model, itemtype='Rasch')
coef(mod, simplify=TRUE)

# Suppose that the first 16 items were suspected to be easier than the last 16 items,
# and we wish to test this item structure hypothesis (more intercept designs are possible
# by including more columns).
itemdesign <- data.frame(itemorder = factor(c(rep('easier', 16), rep('harder', 16))))
rownames(itemdesign) <- colnames(data)
itemdesign

# notice that the 'fixed = ~ ... + items' argument is omitted
LLTM <- mixedmirt(data, model = model, fixed = ~ 0 + itemorder, itemdesign = itemdesign,
  SE = TRUE) # SE argument ensures that the information matrix is computed accurately
summary(LLTM)
coef(LLTM)
wald(LLTM)
L <- matrix(c(-1, 1, 0), 1)
wald(LLTM, L) #first half different from second

# compare to items with estimated slopes (2PL)
twoPL <- mixedmirt(data, model = model, fixed = ~ 0 + itemorder, itemtype = '2PL',
  itemdesign = itemdesign)
# twoPL not mixing too well (AR should be between .2 and .5), decrease MHcand
twoPL <- mixedmirt(data, model = model, fixed = ~ 0 + itemorder, itemtype = '2PL',
  itemdesign = itemdesign, technical = list(MHcand = 0.8))
anova(twoPL, LLTM) #much better fit
summary(twoPL)
coef(twoPL)

wald(twoPL)
L <- matrix(0, 1, 34)
L[1, 1] <- 1
L[1, 2] <- -1

```

```

wald(twoPL, L) # n.s., which is the correct conclusion. Rasch approach gave wrong inference

## LLTM with item error term
LLTMwithError <- mixedmirt(data, model = model, fixed = ~ 0 + itemorder, random = ~ 1|items,
  itemdesign = itemdesign)
summary(LLTMwithError)
# large item level variance after itemorder is regressed; not a great predictor of item difficulty
coef(LLTMwithError)

#####
### Polytomous example

# make an arbitrary group difference
covdat <- data.frame(group = rep(c('m', 'f'), nrow(Science)/2))

# partial credit model
mod <- mixedmirt(Science, covdat, model=1, fixed = ~ 0 + group)
coef(mod)

# gpcm to estimate slopes
mod2 <- mixedmirt(Science, covdat, model=1, fixed = ~ 0 + group,
  itemtype = 'gpcm')
summary(mod2)
anova(mod, mod2)

# graded model
mod3 <- mixedmirt(Science, covdat, model=1, fixed = ~ 0 + group,
  itemtype = 'graded')
coef(mod3)

#####
# latent regression with Rasch and 2PL models

set.seed(1)
n <- 300
a <- matrix(1, 10)
d <- matrix(rnorm(10))
Theta <- matrix(c(rnorm(n, 0), rnorm(n, 1), rnorm(n, 2)))
covdata <- data.frame(group=rep(c('g1','g2','g3'), each=n))
dat <- simdata(a, d, N=n*3, Theta=Theta, itemtype = '2PL')
itemstats(dat)

# had we known the latent abilities, we could have computed the regression coeffs
summary(lm(Theta ~ covdata$group))

# but all we have is observed test data. Latent regression helps to recover these coeffs
# Rasch model approach (and mirt equivalent)
rmod0 <- mirt(dat, 1, 'Rasch') # unconditional

# these two models are equivalent
rmod1a <- mirt(dat, 1, 'Rasch', covdata = covdata, formula = ~ group)
rmod1b <- mixedmirt(dat, covdata, 1, fixed = ~ 0 + items + group)

```

```

anova(rmod0, rmod1b)
coef(rmod1a, simplify=TRUE)
summary(rmod1b)

# 2PL, requires different input to allow Theta variance to remain fixed
mod0 <- mirt(dat, 1) # unconditional
mod1a <- mirt(dat, 1, covdata = covdata, formula = ~ group, itemtype = '2PL')
mod1b <- mixedmirt(dat, covdata, 1, fixed = ~ 0 + items, lr.fixed = ~group, itemtype = '2PL')
anova(mod0, mod1b)
coef(mod1a)$lr.betas
summary(mod1b)

# specifying specific regression effects is accomplished by passing a list of formula
model <- 'F1 = 1-5
        F2 = 6-10'
covdata$contvar <- rnorm(nrow(covdata))
mod2 <- mirt(dat, model, itemtype = 'Rasch', covdata=covdata,
            formula = list(F1 = ~ group + contvar, F2 = ~ group))
coef(mod2)[11:12]
mod2b <- mixedmirt(dat, covdata, model, fixed = ~ 0 + items,
                  lr.fixed = list(F1 = ~ group + contvar, F2 = ~ group))
summary(mod2b)

#####
## Simulated Multilevel Rasch Model

set.seed(1)
N <- 2000
a <- matrix(rep(1,10),10,1)
d <- matrix(rnorm(10))
cluster = 100
random_intercept = rnorm(cluster,0,1)
Theta = numeric()
for (i in 1:cluster)
  Theta <- c(Theta, rnorm(N/cluster,0,1) + random_intercept[i])

group = factor(rep(paste0('G',1:cluster), each = N/cluster))
covdata <- data.frame(group)
dat <- simdata(a,d,N, itemtype = rep('2PL',10), Theta=matrix(Theta))
itemstats(dat)

# null model
mod1 <- mixedmirt(dat, covdata, 1, fixed = ~ 0 + items, random = ~ 1|group)
summary(mod1)

# include level 2 predictor for 'group' variance
covdata$group_pred <- rep(random_intercept, each = N/cluster)
mod2 <- mixedmirt(dat, covdata, 1, fixed = ~ 0 + items + group_pred, random = ~ 1|group)

# including group means predicts nearly all variability in 'group'
summary(mod2)
anova(mod1, mod2)

```

```
# can also be fit for Rasch/non-Rasch models with the lr.random input
mod1b <- mixedmirt(dat, covdata, 1, fixed = ~ 0 + items, lr.random = ~ 1|group)
summary(mod1b)

mod2b <- mixedmirt(dat, covdata, 1, fixed = ~ 0 + items + group_pred, lr.random = ~ 1|group)
summary(mod2b)
anova(mod1b, mod2b)

mod3 <- mixedmirt(dat, covdata, 1, fixed = ~ 0 + items, lr.random = ~ 1|group, itemtype = '2PL')
summary(mod3)
anova(mod1b, mod3)

head(cbind(randef(mod3)$group, random_intercept))

## End(Not run)
```

MixtureClass-class	Class " <i>MixtureClass</i> "
--------------------	-------------------------------

Description

Defines the object returned from `multipleGroup` when estimated with mixture distributions.

Slots

Call: function call

Data: list of data, sometimes in different forms

Options: list of estimation options

Fit: a list of fit information

Model: a list of model-based information

ParObjects: a list of the S4 objects used during estimation

OptimInfo: a list of arguments from the optimization process

Internals: a list of internal arguments for secondary computations (inspecting this object is generally not required)

vcov: a matrix represented the asymptotic covariance matrix of the parameter estimates

time: a data.frame indicating the breakdown of computation times in seconds

Methods

```
coef signature(object = "MixtureClass")
print signature(x = "MixtureClass")
show signature(object = "MixtureClass")
anova signature(object = "MixtureClass")
```

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

mod2values

Convert an estimated mirt model to a data.frame

Description

Given an estimated model from any of mirt's model fitting functions this function will convert the model parameters into the design data frame of starting values and other parameter characteristics (similar to using the pars = 'values' for obtaining starting values).

Usage

```
mod2values(x)
```

Arguments

x an estimated model x from the mirt package

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[extract.mirt](#)

Examples

```
## Not run:
dat <- expand.table(LSAT7)
mod <- mirt(dat, "F=1-5
           CONSTRAIN=(1-5, a1)")
values <- mod2values(mod)
values

# use the converted values as starting values in a new model, and reduce TOL
```

```

mod2 <- mirt(dat, 1, pars = values, TOL=1e-5)
coef(mod2, simplify=TRUE)

# use parameters on different dataset
mod3 <- mirt(expand.table(LSAT6), pars=values)
coef(mod3, simplify=TRUE)

# supports differing itemtypes on second model
sv <- mirt(Science, itemtype=c('graded', rep('gpcm', 3)), pars='values')
mod3 <- mirt(Science, pars = sv) # itemtype omitted
coef(mod3, simplify=TRUE)$items
extract.mirt(mod3, 'itemtype')

## End(Not run)

```

multipleGroup

*Multiple Group Estimation***Description**

multipleGroup performs a full-information maximum-likelihood multiple group analysis for any combination of dichotomous and polytomous data under the item response theory paradigm using either Cai's (2010) Metropolis-Hastings Robbins-Monro (MHRM) algorithm or with an EM algorithm approach. This function may be used for detecting differential item functioning (DIF), though the [DIF](#) function may provide a more convenient approach. If the grouping variable is not specified then the dentype input can be modified to fit mixture models to estimate any latent group components.

Usage

```

multipleGroup(
  data,
  model = 1,
  group,
  itemtype = NULL,
  invariance = "",
  method = "EM",
  dentype = "Gaussian",
  itemdesign = NULL,
  item.formula = NULL,
  ...
)

```

Arguments

data a matrix or data.frame that consists of numerically ordered data, organized in the form of integers, with missing data coded as NA

model	string to be passed to, or a model object returned from, <code>mirt.model</code> declaring how the global model is to be estimated (useful to apply constraints here)
group	a character or factor vector indicating group membership. If a character vector is supplied this will be automatically transformed into a <code>factor</code> variable. As well, the first level of the (factorized) grouping variable will be treated as the "reference" group
itemtype	can be same type of input as is documented in <code>mirt</code> , however may also be a <code>ngroups</code> by <code>nitems</code> matrix specifying the type of IRT models for each group, respectively. Rows of this input correspond to the levels of the group input. For mixture models the rows correspond to the respective mixture grouping variables to be constructed, and the IRT models should be within these mixtures
invariance	a character vector containing the following possible options: 'free_mean' or 'free_means' freely estimate all latent means in all focal groups (reference group constrained to a vector of 0's) 'free_var', 'free_vars', 'free_variance', or 'free_variances' freely estimate all latent variances in focal groups (reference group variances all constrained to 1) 'slopes' to constrain all the slopes to be equal across all groups 'intercepts' to constrain all the intercepts to be equal across all groups, note for nominal models this also includes the category specific slope parameters Additionally, specifying specific item name bundles (from <code>colnames(data)</code>) will constrain all freely estimated parameters in each item to be equal across groups. This is useful for selecting 'anchor' items for vertical and horizontal scaling, and for detecting differential item functioning (DIF) across groups
method	a character object that is either 'EM', 'QMCEM', or 'MHRM' (default is 'EM'). See <code>mirt</code> for details
dentype	type of density form to use for the latent trait parameters. Current options include all of the methods described in <code>mirt</code> , as well as 'mixture-#' estimates mixtures of Gaussian distributions, where the # placeholder represents the number of potential grouping variables (e.g., 'mixture-3' will estimate 3 underlying classes). Each class is assigned the group name MIXTURE_#, where # is the class number. Note that internally the mixture coefficients are stored as log values where the first mixture group coefficient is fixed at 0
itemdesign	see <code>mirt</code> for details
item.formula	see <code>mirt</code> for details
...	additional arguments to be passed to the estimation engine. See <code>mirt</code> for details and examples

Details

By default the estimation in `multipleGroup` assumes that the models are maximally independent, and therefore could initially be performed by sub-setting the data and running identical models with `mirt` and aggregating the results (e.g., log-likelihood). However, constraints may be automatically imposed across groups by invoking various invariance keywords. Users may also supply a list of parameter equality constraints to by `constrain` argument, or define equality constraints using the `mirt.model` syntax (recommended).

Value

function returns an object of class `MultipleGroupClass` ([MultipleGroupClass-class](#)).

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Magnus, B. E. and Garnier-Villarréal (2022). A multidimensional zero-inflated graded response model for ordinal symptom data. *Psychological Methods*, 27, 261-279.

Wall, M., M., Park, J., Y., and Moustaki I. (2015). IRT modeling in the presence of zero-inflation with application to psychiatric disorder severity. *Applied Psychological Measurement* 39: 583-597.

See Also

[mirt](#), [DIF](#), [extract.group](#), [DRF](#)

Examples

```
## Not run:

# single factor
set.seed(12345)
a <- matrix(abs(rnorm(15,1,.3)), ncol=1)
d <- matrix(rnorm(15,0,.7),ncol=1)
itemtype <- rep('2PL', nrow(a))
N <- 1000
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a, d, N, itemtype, mu = .1, sigma = matrix(1.5))
dat <- rbind(dataset1, dataset2)
group <- c(rep('D1', N), rep('D2', N))

# marginal information
itemstats(dat)

# conditional information
itemstats(dat, group=group)

mod_configural <- multipleGroup(dat, 1, group = group) #completely separate analyses
# limited information fit statistics
M2(mod_configural)

mod_metric <- multipleGroup(dat, 1, group = group, invariance=c('slopes')) #equal slopes
# equal intercepts, free variance and means
mod_scalar2 <- multipleGroup(dat, 1, group = group,
                             invariance=c('slopes', 'intercepts', 'free_var','free_means'))
mod_scalar1 <- multipleGroup(dat, 1, group = group, #fixed means
                             invariance=c('slopes', 'intercepts', 'free_var'))
```

```

mod_fullconstrain <- multipleGroup(dat, 1, group = group,
                                   invariance=c('slopes', 'intercepts'))
extract.mirt(mod_fullconstrain, 'time') #time of estimation components

# optionally use Newton-Raphson for (generally) faster convergence in the
# M-step's, though occasionally less stable
mod_fullconstrain <- multipleGroup(dat, 1, group = group, optimizer = 'NR',
                                   invariance=c('slopes', 'intercepts'))
extract.mirt(mod_fullconstrain, 'time') #time of estimation components

summary(mod_scalar2)
coef(mod_scalar2, simplify=TRUE)
residuals(mod_scalar2)
plot(mod_configural)
plot(mod_configural, type = 'info')
plot(mod_configural, type = 'trace')
plot(mod_configural, type = 'trace', which.items = 1:4)
itemplot(mod_configural, 2)
itemplot(mod_configural, 2, type = 'RE')

anova(mod_metric, mod_configural) #equal slopes only
anova(mod_scalar2, mod_metric) #equal intercepts, free variance and mean
anova(mod_scalar1, mod_scalar2) #fix mean
anova(mod_fullconstrain, mod_scalar1) #fix variance

# compared all at once (in order of most constrained to least)
anova(mod_fullconstrain, mod_scalar2, mod_configural)

# test whether first 6 slopes should be equal across groups
values <- multipleGroup(dat, 1, group = group, pars = 'values')
values
constrain <- list(c(1, 63), c(5,67), c(9,71), c(13,75), c(17,79), c(21,83))
equalslopes <- multipleGroup(dat, 1, group = group, constrain = constrain)
anova(equalslopes, mod_configural)

# same as above, but using mirt.model syntax
newmodel <- '
  F = 1-15
  CONSTRAINB = (1-6, a1)'
equalslopes <- multipleGroup(dat, newmodel, group = group)
coef(equalslopes, simplify=TRUE)

#####
# vertical scaling (i.e., equating when groups answer items others do not)
dat2 <- dat
dat2[group == 'D1', 1:2] <- dat2[group != 'D1', 14:15] <- NA
head(dat2)
tail(dat2)

# items with missing responses need to be constrained across groups for identification
nms <- colnames(dat2)
mod <- multipleGroup(dat2, 1, group, invariance = nms[c(1:2, 14:15)])

```

```

# this will throw an error without proper constraints (SEs cannot be computed either)
# mod <- multipleGroup(dat2, 1, group)

# model still does not have anchors, therefore need to add a few (here use items 3-5)
mod_anchor <- multipleGroup(dat2, 1, group,
                             invariance = c(nms[c(1:5, 14:15)], 'free_means', 'free_var'))
coef(mod_anchor, simplify=TRUE)

# check if identified by computing information matrix
mod_anchor <- multipleGroup(dat2, 1, group, pars = mod2values(mod_anchor), TOL=NaN, SE=TRUE,
                             invariance = c(nms[c(1:5, 14:15)], 'free_means', 'free_var'))
mod_anchor
coef(mod_anchor)
coef(mod_anchor, printSE=TRUE)

#####
# DIF test for each item (using all other items as anchors)
itemnames <- colnames(dat)
refmodel <- multipleGroup(dat, 1, group = group, SE=TRUE,
                           invariance=c('free_means', 'free_var', itemnames))

# loop over items (in practice, run in parallel to increase speed). May be better to use ?DIF
estmodels <- vector('list', ncol(dat))
for(i in 1:ncol(dat))
  estmodels[[i]] <- multipleGroup(dat, 1, group = group, verbose = FALSE,
                                  invariance=c('free_means', 'free_var', itemnames[-i]))
anova(refmodel, estmodels[[1]])
(anovas <- lapply(estmodels, function(x, refmodel) anova(refmodel, x),
                  refmodel=refmodel))

# family-wise error control
p <- do.call(rbind, lapply(anovas, function(x) x[2, 'p']))
p.adjust(p, method = 'BH')

# same as above, except only test if slopes vary (1 df)
# constrain all intercepts
estmodels <- vector('list', ncol(dat))
for(i in 1:ncol(dat))
  estmodels[[i]] <- multipleGroup(dat, 1, group = group, verbose = FALSE,
                                  invariance=c('free_means', 'free_var', 'intercepts',
                                                itemnames[-i]))

(anovas <- lapply(estmodels, function(x, refmodel) anova(refmodel, x),
                  refmodel=refmodel))

# quickly test with Wald test using DIF()
mod_configural2 <- multipleGroup(dat, 1, group = group, SE=TRUE)
DIF(mod_configural2, which.par = c('a1', 'd'), Wald=TRUE, p.adjust = 'fdr')

```

```
#####
# Three group model where the latent variable parameters are constrained to
# be equal in the focal groups

set.seed(12345)
a <- matrix(abs(rnorm(15,1,.3)), ncol=1)
d <- matrix(rnorm(15,0,.7),ncol=1)
itemtype <- rep('2PL', nrow(a))
N <- 1000
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a, d, N, itemtype, mu = .1, sigma = matrix(1.5))
dataset3 <- simdata(a, d, N, itemtype, mu = .1, sigma = matrix(1.5))
dat <- rbind(dataset1, dataset2, dataset3)
group <- rep(c('D1', 'D2', 'D3'), each=N)

# marginal information
itemstats(dat)

# conditional information
itemstats(dat, group=group)

model <- 'F1 = 1-15
          FREE[D2, D3] = (GROUP, MEAN_1), (GROUP, COV_11)
          CONSTRAINB[D2,D3] = (GROUP, MEAN_1), (GROUP, COV_11)'

mod <- multipleGroup(dat, model, group = group, invariance = colnames(dat))
coef(mod, simplify=TRUE)

#####
# Testing main effects in multiple independent grouping variables
set.seed(1234)
a <- matrix(abs(rnorm(15,1,.3)), ncol=1)
d <- matrix(rnorm(15,0,.7),ncol=1)
itemtype <- rep('2PL', nrow(a))
N <- 500

# generated data have interaction effect for latent means, as well as a
# main effect across D but no main effect across G
d11 <- simdata(a, d, N, itemtype, mu = 0)
d12 <- simdata(a, d, N, itemtype, mu = 0)
d13 <- simdata(a, d, N, itemtype, mu = 0)
d21 <- simdata(a, d, N, itemtype, mu = 1/2)
d22 <- simdata(a, d, N, itemtype, mu = 1/2)
d23 <- simdata(a, d, N, itemtype, mu = -1)
dat <- do.call(rbind, list(d11, d12, d13, d21, d22, d23))
group <- rep(c('G1.D1', 'G1.D2', 'G1.D3', 'G2.D1', 'G2.D2', 'G2.D3'), each=N)
table(group)

# in practice, group would be organized in a data.frame as follows
df <- data.frame(group)
dfw <- tidyr::separate_wider_delim(df, group, delim='.', names = c('G', 'D'))
head(dfw)
```

```

# for use with multipleGroup() combine into a single long group
group <- with(dfw, factor(G):factor(D))

# conditional information
itemstats(dat, group=group)

mod <- multipleGroup(dat, group = group, SE=TRUE,
                     invariance = c(colnames(dat), 'free_mean', 'free_var'))
coef(mod, simplify=TRUE)
sapply(coef(mod, simplify=TRUE), \(x) unname(x$means)) # mean estimates
wald(mod) # view parameter names for later testing

# test for main effect over G group (manually compute marginal mean)
wald(mod, "0 + MEAN_1.123 + MEAN_1.185 = MEAN_1.247 + MEAN_1.309 + MEAN_1.371")

# test for main effect over D group (manually compute marginal means)
wald(mod, c("0 + MEAN_1.247 = MEAN_1.123 + MEAN_1.309",
            "0 + MEAN_1.247 = MEAN_1.185 + MEAN_1.371"))

# post-hoc tests (better practice would include p.adjust() )
wald(mod, "0 + MEAN_1.247 = MEAN_1.123 + MEAN_1.309") # D1 vs D2
wald(mod, "0 + MEAN_1.247 = MEAN_1.185 + MEAN_1.371") # D1 vs D3
wald(mod, "MEAN_1.123 + MEAN_1.309 = MEAN_1.185 + MEAN_1.371") # D2 vs D3

#####
# multiple factors

a <- matrix(c(abs(rnorm(5,1,.3)), rep(0,15),abs(rnorm(5,1,.3)),
              rep(0,15),abs(rnorm(5,1,.3))), 15, 3)
d <- matrix(rnorm(15,0,.7),ncol=1)
mu <- c(-.4, -.7, .1)
sigma <- matrix(c(1.21,.297,1.232,.297,.81,.252,1.232,.252,1.96),3,3)
itemtype <- rep('2PL', nrow(a))
N <- 1000
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a, d, N, itemtype, mu = mu, sigma = sigma)
dat <- rbind(dataset1, dataset2)
group <- c(rep('D1', N), rep('D2', N))

# group models
model <- '
  F1 = 1-5
  F2 = 6-10
  F3 = 11-15'

# define mirt cluster to use parallel architecture
if(interactive()) mirtCluster()

# EM approach (not as accurate with 3 factors, but generally good for quick model comparisons)
mod_configural <- multipleGroup(dat, model, group = group) #completely separate analyses
mod_metric <- multipleGroup(dat, model, group = group, invariance=c('slopes')) #equal slopes
mod_fullconstrain <- multipleGroup(dat, model, group = group, #equal means, slopes, intercepts

```

```

                                invariance=c('slopes', 'intercepts'))

anova(mod_metric, mod_configural)
anova(mod_fullconstrain, mod_metric)

# same as above, but with MHRM (generally more accurate with 3+ factors, but slower)
mod_configural <- multipleGroup(dat, model, group = group, method = 'MHRM')
mod_metric <- multipleGroup(dat, model, group = group, invariance=c('slopes'), method = 'MHRM')
mod_fullconstrain <- multipleGroup(dat, model, group = group, method = 'MHRM',
                                   invariance=c('slopes', 'intercepts'))

anova(mod_metric, mod_configural)
anova(mod_fullconstrain, mod_metric)

#####
# polytomous item example
set.seed(12345)
a <- matrix(abs(rnorm(15,1,.3)), ncol=1)
d <- matrix(rnorm(15,0,.7), ncol=1)
d <- cbind(d, d-1, d-2)
itemtype <- rep('graded', nrow(a))
N <- 1000
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a, d, N, itemtype, mu = .1, sigma = matrix(1.5))
dat <- rbind(dataset1, dataset2)
group <- c(rep('D1', N), rep('D2', N))
model <- 'F1 = 1-15'

mod_configural <- multipleGroup(dat, model, group = group)
plot(mod_configural)
plot(mod_configural, type = 'SE')
itemplot(mod_configural, 1)
itemplot(mod_configural, 1, type = 'info')
plot(mod_configural, type = 'trace') # messy, score function typically better
plot(mod_configural, type = 'itemscore')

fs <- fscores(mod_configural, full.scores = FALSE)
head(fs[["D1"]])
fscores(mod_configural, method = 'EAPsum', full.scores = FALSE)

# constrain slopes within each group to be equal (but not across groups)
model2 <- 'F1 = 1-15
          CONSTRAIN = (1-15, a1)'
mod_configural2 <- multipleGroup(dat, model2, group = group)
plot(mod_configural2, type = 'SE')
plot(mod_configural2, type = 'RE')
itemplot(mod_configural2, 10)

#####
## empirical histogram example (normal and bimodal groups)
set.seed(1234)
a <- matrix(rlnorm(50, .2, .2))
d <- matrix(rnorm(50))

```

```

ThetaNormal <- matrix(rnorm(2000))
ThetaBimodal <- scale(matrix(c(rnorm(1000, -2), rnorm(1000, 2)))) #bimodal
Theta <- rbind(ThetaNormal, ThetaBimodal)
dat <- simdata(a, d, 4000, itemtype = '2PL', Theta=Theta)
group <- rep(c('G1', 'G2'), each=2000)

EH <- multipleGroup(dat, 1, group=group, dentype="empiricalhist", invariance = colnames(dat))
coef(EH, simplify=TRUE)
plot(EH, type = 'empiricalhist', npts = 60)

# DIF test for item 1
EH1 <- multipleGroup(dat, 1, group=group, dentype="empiricalhist", invariance = colnames(dat)[-1])
anova(EH, EH1)

#-----
# Mixture model (no prior group variable specified)

set.seed(12345)
nitems <- 20
a1 <- matrix(.75, ncol=1, nrow=nitems)
a2 <- matrix(1.25, ncol=1, nrow=nitems)
d1 <- matrix(rnorm(nitems,0,1),ncol=1)
d2 <- matrix(rnorm(nitems,0,1),ncol=1)
itemtype <- rep('2PL', nrow(a1))
N1 <- 500
N2 <- N1*2 # second class twice as large

dataset1 <- simdata(a1, d1, N1, itemtype)
dataset2 <- simdata(a2, d2, N2, itemtype)
dat <- rbind(dataset1, dataset2)
# group <- c(rep('D1', N1), rep('D2', N2))

# Mixture Rasch model (Rost, 1990)
models <- 'F1 = 1-20
          CONSTRAIN = (1-20, a1)'
mod_mix <- multipleGroup(dat, models, dentype = 'mixture-2', GenRandomPars = TRUE)
coef(mod_mix, simplify=TRUE)
summary(mod_mix)
plot(mod_mix)
plot(mod_mix, type = 'trace')
itemplot(mod_mix, 1, type = 'info')

head(fscores(mod_mix)) # theta estimates
head(fscores(mod_mix, method = 'classify')) # classification probability
itemfit(mod_mix)

# Mixture 2PL model
mod_mix2 <- multipleGroup(dat, 1, dentype = 'mixture-2', GenRandomPars = TRUE)
anova(mod_mix, mod_mix2)
coef(mod_mix2, simplify=TRUE)
itemfit(mod_mix2)

# Compare to single group

```

```

mod <- mirt(dat)
anova(mod, mod_mix2)

#####
# Zero-inflated 2PL IRT model (Wall, Park, and Moustaki, 2015)

n <- 1000
nitems <- 20

a <- rep(2, nitems)
d <- rep(c(-2,-1,0,1,2), each=nitems/5)
zi_p <- 0.2 # Proportion of people in zero class

theta <- rnorm(n, 0, 1)
zeros <- matrix(0, n*zi_p, nitems)
nonzeros <- simdata(a, d, n*(1-zi_p), itemtype = '2PL',
  Theta = as.matrix(theta[1:(n*(1-zi_p))]))
data <- rbind(nonzeros, zeros)

# define class with extreme theta but fixed item parameters
zi2PL <- "F = 1-20
  START [MIXTURE_1] = (GROUP, MEAN_1, -100), (GROUP, COV_11, .00001),
    (1-20, a1, 1.0), (1-20, d, 0)
  FIXED [MIXTURE_1] = (GROUP, MEAN_1), (GROUP, COV_11),
    (1-20, a1), (1-20, d)"

# define custom Theta integration grid that contains extreme theta + normal grid
technical <- list(customTheta = matrix(c(-100, seq(-6,6,length.out=61))))

# fit ZIM-IRT
zi2PL.fit <- multipleGroup(data, zi2PL, dentype = 'mixture-2', technical=technical)
coef(zi2PL.fit, simplify=TRUE)

# classification estimates
pi_hat <- fscores(zi2PL.fit, method = 'classify')
head(pi_hat)
tail(pi_hat)

# EAP estimates (not useful for zip class)
fs <- fscores(zi2PL.fit)
head(fs)
tail(fs)

#####
# Zero-inflated graded response model (Magnus and Garnier-Villarre, 2022)

n <- 1000
nitems <- 20

a <- matrix(rlnorm(20,.2,.3))

# for the graded model, ensure that there is enough space between the intercepts,
# otherwise closer categories will not be selected often (minimum distance of 0.3 here)

```

```

diffs <- t(apply(matrix(runif(20*4, .3, 1), 20), 1, cumsum))
diffs <- -(diffs - rowMeans(diffs))
d <- diffs + rnorm(20)

zi_p <- 0.2 # Proportion of people in zero/lowest category class

theta <- rnorm(n, 0, 1)
zeros <- matrix(0, n*zi_p, nitems)
nonzeros <- simdata(a, d, n*(1-zi_p), itemtype = 'graded',
  Theta = as.matrix(theta[1:(n*(1-zi_p))]))
data <- rbind(nonzeros, zeros)

# intercepts will be labelled as d1 through d4
apply(data, 2, table)

# ignoring zero inflation (bad idea)
modGRM <- mirt(data)
coef(modGRM, simplify=TRUE)

# Define class with extreme theta but fixed item parameters
# For GRM in zero-inflated class the intercept values are arbitrary
# as the model forces the responses all into the first category (hence,
# spacing arbitrarily set to 1)
ziGRM <- "F = 1-20
  START [MIXTURE_1] = (GROUP, MEAN_1, -100), (GROUP, COV_11, .00001),
    (1-20, a1, 1.0),
    (1-20, d1, 2), (1-20, d2, 1), (1-20, d3, 0), (1-20, d4, -1)
  FIXED [MIXTURE_1] = (GROUP, MEAN_1), (GROUP, COV_11),
    (1-20, a1),
    (1-20, d1), (1-20, d2), (1-20, d3), (1-20, d4)"

# define custom Theta integration grid that contains extreme theta + normal grid
technical <- list(customTheta = matrix(c(-100, seq(-6,6,length.out=61))))

# fit zero-inflated GRM
ziGRM.fit <- multipleGroup(data, ziGRM, dentype = 'mixture-2', technical=technical)
coef(ziGRM.fit, simplify=TRUE)

# classification estimates
pi_hat <- fscores(ziGRM.fit, method = 'classify')
head(pi_hat)
tail(pi_hat)

# EAP estimates (not useful for zip class)
fs <- fscores(ziGRM.fit)
head(fs)
tail(fs)

## End(Not run)

```

MultipleGroupClass-class
Class "MultipleGroupClass"

Description

Defines the object returned from `multipleGroup`.

Slots

Call: function call

Data: list of data, sometimes in different forms

Options: list of estimation options

Fit: a list of fit information

Model: a list of model-based information

ParObjects: a list of the S4 objects used during estimation

OptimInfo: a list of arguments from the optimization process

Internals: a list of internal arguments for secondary computations (inspecting this object is generally not required)

vcov: a matrix represented the asymptotic covariance matrix of the parameter estimates

time: a data.frame indicating the breakdown of computation times in seconds

Methods

coef signature(object = "MultipleGroupClass")

print signature(x = "MultipleGroupClass")

show signature(object = "MultipleGroupClass")

anova signature(object = "MultipleGroupClass")

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

numerical_deriv	<i>Compute numerical derivatives</i>
-----------------	--------------------------------------

Description

Compute numerical derivatives using forward/backward difference, central difference, or Richardson extrapolation.

Usage

```
numerical_deriv(
  par,
  f,
  ...,
  delta = 1e-05,
  gradient = TRUE,
  type = "Richardson"
)
```

Arguments

par	a vector of parameters to find partial derivative at
f	the objective function being evaluated
...	additional arguments to be passed to f
delta	the term used to perturb the f function. Default is 1e-5
gradient	logical; compute the gradient terms? If FALSE then the Hessian is computed instead
type	type of difference to compute. Can be either 'forward' for the forward difference, 'central' for the central difference, or 'Richardson' for the Richardson extrapolation (default). Backward difference is achieved by supplying a negative delta value with 'forward'. When type = 'Richardson', the default value of delta is increased to $\delta * 1000$ for the Hessian and $\delta * 10$ for the gradient to provide a reasonable perturbation starting location (each delta is halved at each iteration).

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

Examples

```
## Not run:
f <- function(x) 3*x[1]^3 - 4*x[2]^2
par <- c(3,8)

# grad = 9 * x^2 , -8 * y
```

```

(actual <- c(9 * par[1]^2, -8 * par[2]))
numerical_deriv(par, f, type = 'forward')
numerical_deriv(par, f, type = 'central')
numerical_deriv(par, f, type = 'Richardson') # default

# Hessian = h11 -> 18 * x, h22 -> -8, h12 -> h21 -> 0
(actual <- matrix(c(18 * par[1], 0, 0, -8), 2, 2))
numerical_deriv(par, f, type = 'forward', gradient = FALSE)
numerical_deriv(par, f, type = 'central', gradient = FALSE)
numerical_deriv(par, f, type = 'Richardson', gradient = FALSE) # default

## End(Not run)

```

personfit

Person fit statistics

Description

personfit calculates the Zh values from Drasgow, Levine and Williams (1985) for unidimensional and multidimensional models, as well as the infit and outfit statistics. The returned object is a data.frame consisting either of the tabulated data or full data with the statistics appended to the rightmost columns.

Usage

```

personfit(
  x,
  method = "EAP",
  Theta = NULL,
  stats.only = TRUE,
  return.resids = FALSE,
  ...
)

```

Arguments

x	a computed model object of class SingleGroupClass or MultipleGroupClass
method	type of factor score estimation method. See fscores for more detail
Theta	a matrix of factor scores used for statistics that require empirical estimates. If supplied, arguments typically passed to fscores() will be ignored and these values will be used instead
stats.only	logical; return only the person fit statistics without their associated response pattern?
return.resids	logical; return the standardized and unstandardized N by J matrices of person and item residuals? If TRUE will return a named list of each residual type
...	additional arguments to be passed to fscores()

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Drasgow, F., Levine, M. V., & Williams, E. A. (1985). Appropriateness measurement with polychotomous item response models and standardized indices. *British Journal of Mathematical and Statistical Psychology*, 38, 67-86.
- Reise, S. P. (1990). A comparison of item- and person-fit methods of assessing model-data fit in IRT. *Applied Psychological Measurement*, 14, 127-137.
- Wright B. D. & Masters, G. N. (1982). *Rating scale analysis*. MESA Press.

See Also

[itemfit](#)

Examples

```
## Not run:

#make some data
set.seed(1)
a <- matrix(rlnorm(20),ncol=1)
d <- matrix(rnorm(20),ncol=1)
items <- rep('2PL', 20)
data <- simdata(a,d, 2000, items)

# first observation responds 1 for most difficult, 0 for easiest
data[1,] <- ifelse(d > 0, 0, 1)

# second observations answers first half as 1 second half as 0
data[2,] <- rep(1:0, each = 10)

x <- mirt(data, 1)
fit <- personfit(x)
head(fit)

# raw/standardized residuals
resid_list <- personfit(x, return.resids=TRUE)
head(resid_list$resid) # unstandardized
head(resid_list$std.resid) # standardized (approximate z-scores)

# with missing data
data[3, c(1,3,5,7)] <- NA
x.miss <- mirt(data, 1)
fit.miss <- personfit(x.miss)
head(fit.miss)
head(personfit(x.miss, return.resids=TRUE))
```

```

#using precomputed Theta
Theta <- fscores(x, method = 'MAP', full.scores = TRUE)
head(personfit(x, Theta=Theta))

# multiple group Rasch model example
set.seed(12345)
a <- matrix(rep(1, 16), ncol=1)
d <- matrix(rnorm(16,0,.7),ncol=1)
itemtype <- rep('dich', nrow(a))
N <- 1000
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a, d, N, itemtype, sigma = matrix(1.5))
dat <- rbind(dataset1, dataset2)

# first observation responds 1 for most difficult, 0 for easiest
dat[1,] <- ifelse(d > 0, 0, 1)

group <- c(rep('D1', N), rep('D2', N))
models <- 'F1 = 1-16'
mod_Rasch <- multipleGroup(dat, models, itemtype = 'Rasch', group = group)
coef(mod_Rasch, simplify=TRUE)
pf <- personfit(mod_Rasch, method='MAP')
head(pf)

## End(Not run)

```

PLCI.mirt

Compute profiled-likelihood (or posterior) confidence intervals

Description

Computes profiled-likelihood based confidence intervals. Supports the inclusion of equality constraints. Object returns the confidence intervals and whether the respective interval could be found.

Usage

```

PLCI.mirt(
  mod,
  parnum = NULL,
  alpha = 0.05,
  search_bound = TRUE,
  step = 0.5,
  lower = TRUE,
  upper = TRUE,
  inf2val = 30,
  NealeMiller = FALSE,

```

```

    verbose = TRUE,
    ...
  )

```

Arguments

<code>mod</code>	a converged mirt model
<code>parnum</code>	a numeric vector indicating which parameters to estimate. Use <code>mod2values</code> to determine parameter numbers. If NULL, all possible parameters are used
<code>alpha</code>	two-tailed alpha critical level
<code>search_bound</code>	logical; use a fixed grid of values around the ML estimate to determine more suitable optimization bounds? Using this has much better behaviour than setting fixed upper/lower bound values and searching from more extreme ends
<code>step</code>	magnitude of steps used when <code>search_bound</code> is TRUE. Smaller values create more points to search a suitable bound for (up to the lower bound value visible with <code>mod2values</code>). When upper/lower bounds are detected this value will be adjusted accordingly
<code>lower</code>	logical; search for the lower CI?
<code>upper</code>	logical; search for the upper CI?
<code>inf2val</code>	a numeric used to change parameter bounds which are infinity to a finite number. Decreasing this too much may not allow a suitable bound to be located. Default is 30
<code>NealeMiller</code>	logical; use the Neale and Miller 1997 approximation? Default is FALSE
<code>verbose</code>	logical; include additional information in the console?
<code>...</code>	additional arguments to pass to the estimation functions

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Chalmers, R. P., Pek, J., & Liu, Y. (2017). Profile-likelihood Confidence Intervals in Item Response Theory Models. *Multivariate Behavioral Research*, 52, 533-550. doi:10.1080/00273171.2017.1329082
- Neale, M. C. & Miller, M. B. (1997). The use of likelihood-based confidence intervals in genetic models. *Behavior Genetics*, 27, 113-120.

See Also

[boot.mirt](#)

Examples

```
## Not run:
if(interactive()) mirtCluster() #use all available cores to estimate CI's in parallel
dat <- expand.table(LSAT7)
mod <- mirt(dat, 1)

result <- PLCI.mirt(mod)
result

# model with constraints
mod <- mirt(dat, 'F = 1-5
              CONSTRAIN = (1-5, a1)')

result <- PLCI.mirt(mod)
result

mod2 <- mirt(Science, 1)
result2 <- PLCI.mirt(mod2)
result2

# only estimate CI's slopes
sv <- mod2values(mod2)
parnum <- sv$parnum[sv$name == 'a1']
result3 <- PLCI.mirt(mod2, parnum)
result3

## End(Not run)
```

plot, MultipleGroupClass, missing-method

Plot various test-implied functions from models

Description

Plot various test implied response functions from models estimated in the mirt package.

Usage

```
## S4 method for signature 'MultipleGroupClass,missing'
plot(
  x,
  y,
  type = "score",
  npts = 200,
  drop2 = TRUE,
  degrees = 45,
  which.items = 1:extract.mirt(x, "nitems"),
  rot = list(xaxis = -70, yaxis = 30, zaxis = 10),
```

```

    facet_items = TRUE,
    theta_lim = c(-6, 6),
    par.strip.text = list(cex = 0.7),
    par.settings = list(strip.background = list(col = "#9ECAE1"), strip.border = list(col =
      "black")),
    auto.key = list(space = "right", points = FALSE, lines = TRUE),
    ...
  )

## S4 method for signature 'SingleGroupClass,missing'
plot(
  x,
  y,
  type = "score",
  npts = 200,
  drop2 = TRUE,
  degrees = 45,
  theta_lim = c(-6, 6),
  which.items = 1:extract.mirt(x, "nitems"),
  MI = 0,
  CI = 0.95,
  rot = list(xaxis = -70, yaxis = 30, zaxis = 10),
  facet_items = TRUE,
  main = NULL,
  drape = TRUE,
  colorkey = TRUE,
  ehist.cut = 1e-10,
  add.ylab2 = TRUE,
  par.strip.text = list(cex = 0.7),
  par.settings = list(strip.background = list(col = "#9ECAE1"), strip.border = list(col =
    "black")),
  auto.key = list(space = "right", points = FALSE, lines = TRUE),
  profile = FALSE,
  ...
)

```

Arguments

x	an object of class SingleGroupClass, MultipleGroupClass, or DiscreteClass
y	an arbitrary missing argument required for R CMD check
type	type of plot to view. Can be 'info' test information function 'rxx' for the reliability function 'infocontour' for the test information contours 'SE' for the test standard error function 'infotrace' item information traceline plots 'infoSE' a combined test information and standard error plot

	'trace' item probability traceline plots
	'itemscore' item scoring traceline plots
	'score' expected total score surface
	'scorecontour' expected total score contour plot
	'posteriorTheta' posterior for the latent trait distribution
	'EAPsum' compares sum-scores to the expected values based on the EAP for sum-scores method (see fscores)
	Note that if <code>dentype = 'empiricalhist'</code> was used in estimation then the type 'empiricalhist' also will be available to generate the empirical histogram plot, and if <code>dentype = 'Davidian-#'</code> was used then the type 'Davidian' will also be available to generate the curve estimates at the quadrature nodes used during estimation
npts	number of quadrature points to be used for plotting features. Larger values make plots look smoother
drop2	logical; where appropriate, for dichotomous response items drop the lowest category and provide information pertaining only to the second response option?
degrees	numeric value ranging from 0 to 90 used in plot to compute angle for information-based plots with respect to the first dimension. If a vector is used then a bubble plot is created with the summed information across the angles specified (e.g., <code>degrees = seq(0, 90, by=10)</code>)
which.items	numeric vector indicating which items to be used when plotting. Default is to use all available items
rot	allows rotation of the 3D graphics
facet.items	logical; apply grid of plots across items? If FALSE, items will be placed in one plot for each group
theta_lim	lower and upper limits of the latent trait (theta) to be evaluated, and is used in conjunction with npts
par.strip.text	plotting argument passed to lattice
par.settings	plotting argument passed to lattice
auto.key	plotting argument passed to lattice
...	additional arguments to be passed to lattice
MI	a single number indicating how many imputations to draw to form bootstrapped confidence intervals for the selected test statistic. If greater than 0 a plot will be drawn with a shaded region for the interval
CI	a number from 0 to 1 indicating the confidence interval to select when MI input is used. Default uses the 95% confidence (<code>CI = .95</code>)
main	argument passed to lattice. Default generated automatically
drape	logical argument passed to lattice. Default generated automatically
colorkey	logical argument passed to lattice. Default generated automatically
ehist.cut	a probability value indicating a threshold for excluding cases in empirical histogram plots. Values larger than the default will include more points in the tails of the plot, potentially squishing the 'meat' of the plot to take up less area than visually desired

add.ylab2	logical argument passed to lattice. Default generated automatically
profile	logical; provide a profile plot of response probabilities (objects returned from mdirt only)

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```
## Not run:
x <- mirt(Science, 1, SE=TRUE)
plot(x)
plot(x, type = 'info')
plot(x, type = 'infotrace')
plot(x, type = 'infotrace', facet_items = FALSE)
plot(x, type = 'infoSE')
plot(x, type = 'rxx')
plot(x, type = 'posteriorTheta')

# confidence interval plots when information matrix computed
plot(x)
plot(x, MI=100)
plot(x, type='info', MI=100)
plot(x, type='SE', MI=100)
plot(x, type='rxx', MI=100)

# use the directlabels package to put labels on tracelines
library(directlabels)
plt <- plot(x, type = 'trace')
direct.label(plt, 'top.points')

# additional modifications can be made via update().
# See ?update.trellis for further documentation
plt
update(plt, ylab = expression(Prob(theta)),
      main = "Item Traceline Functions") # ylab/main changed

set.seed(1234)
group <- sample(c('g1','g2'), nrow(Science), TRUE)
x2 <- multipleGroup(Science, 1, group)
plot(x2)
plot(x2, type = 'trace')
plot(x2, type = 'trace', which.items = 1:2)
plot(x2, type = 'itemscore', which.items = 1:2)
plot(x2, type = 'trace', which.items = 1, facet_items = FALSE) #facet by group
plot(x2, type = 'info')

x3 <- mirt(Science, 2)
plot(x3, type = 'info')
plot(x3, type = 'SE', theta_lim = c(-3,3))
```

```
## End(Not run)
```

poly2dich	<i>Change polytomous items to dichotomous item format</i>
-----------	---

Description

Transforms a matrix of items into a new matrix where the select polytomous items have been converted into comparable dichotomous items with the same information.

Usage

```
poly2dich(data, which.items = 1:ncol(data), sep = "_cat.")
```

Arguments

data	an object of class <code>data.frame</code> or <code>matrix</code>
which.items	a vector indicating which items should be transformed into the dichotomous form. Default uses all input items
sep	character vector pattern to append to each item name in data

Value

Returns an integer matrix

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
## Not run:
data(Science)

head(Science)
newScience <- poly2dich(Science)
head(newScience)

newScience2 <- poly2dich(Science, which.items = 2)
head(newScience2)
```

```
## End(Not run)
```

print-method	<i>Print the model objects</i>
--------------	--------------------------------

Description

Print model object summaries to the console.

Usage

```
## S4 method for signature 'SingleGroupClass'
print(x)
```

Arguments

x an object of class SingleGroupClass, MultipleGroupClass, or MixedClass

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```
## Not run:
x <- mirt(Science, 1)
print(x)

## End(Not run)
```

print.mirt_df	<i>Print generic for customized data.frame console output</i>
---------------	---

Description

Provides a nicer output for most printed data.frame objects defined by functions in mirt.

Usage

```
## S3 method for class 'mirt_df'
print(x, digits = 3, ...)
```

Arguments

x	object of class 'mirt_df'
digits	number of digits to round
...	additional arguments passed to print(...)

print.mirt_list	<i>Print generic for customized list console output</i>
-----------------	---

Description

Provides a nicer output for most printed list objects defined by functions in mirt.

Usage

```
## S3 method for class 'mirt_list'  
print(x, digits = 3, ...)
```

Arguments

x	object of class 'mirt_list'
digits	number of digits to round
...	additional arguments passed to print(...)

print.mirt_matrix	<i>Print generic for customized matrix console output</i>
-------------------	---

Description

Provides a nicer output for most printed matrix objects defined by functions in mirt.

Usage

```
## S3 method for class 'mirt_matrix'  
print(x, digits = 3, ...)
```

Arguments

x	object of class 'mirt_matrix'
digits	number of digits to round
...	additional arguments passed to print(...)

probtrace	<i>Function to calculate probability trace lines</i>
-----------	--

Description

Given an internal mirt object extracted from an estimated model, or the single-group estimated model itself, compute the probability trace lines for all categories.

Usage

```
probtrace(x, Theta)
```

Arguments

x	either an extracted internal mirt object containing item information (see extract.item) or a model of class SingleGroupClass typically returned by the function mirt or bfactor
Theta	a vector (unidimensional) or matrix (unidimensional/multidimensional) of latent trait values

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[extract.item](#), [extract.group](#)

Examples

```
mod <- mirt(Science, 1)

# single item probability tracelines for Item 2
extr.2 <- extract.item(mod, 2)
Theta <- matrix(seq(-4,4, by = .1))
traceline <- probtrace(extr.2, Theta)
head(data.frame(traceline, Theta=Theta))

# probability tracelines for all items in test
tracelines <- probtrace(mod, Theta)
head(tracelines)
```

randef	<i>Compute posterior estimates of random effect</i>
--------	---

Description

Stochastically compute random effects for MixedClass objects with Metropolis-Hastings samplers and averaging over the draws to obtain expected a posteriori predictions. Returns a list of the estimated effects.

Usage

```
randef(x, ndraws = 1000, thin = 10, return.draws = FALSE)
```

Arguments

x	an estimated model object from the <code>mixedmirt</code> function
ndraws	total number of draws to perform. Default is 1000
thin	amount of thinning to apply. Default is to use every 10th draw
return.draws	logical; return a list containing the thinned draws of the posterior?

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29.

Chalmers, R. P. (2015). Extended Mixed-Effects Item Response Models with the MH-RM Algorithm. *Journal of Educational Measurement*, 52, 200-222. doi:10.1111/jedm.12072 doi:10.18637/jss.v048.i06

Examples

```
## Not run:
# make an arbitrary groups
covdat <- data.frame(group = rep(paste0('group', 1:49), each=nrow(Science)/49))

# partial credit model
mod <- mixedmirt(Science, covdat, model=1, random = ~ 1|group)
summary(mod)

effects <- randef(mod, ndraws = 2000, thin = 20)
head(effects$Theta)
head(effects$group)

## End(Not run)
```

Description

Computes an IRT version of the "reliable change index" (RCI) proposed by Jacobson and Traux (1991) but modified to use IRT information about scores and measurement error (see Jabrayilov, Emons, and Sijtsma (2016)). Main benefit of the IRT approach is the inclusion of response pattern information in the pre/post data score estimates, as well as conditional standard error of measurement information.

Usage

```
RCI(
  mod_pre,
  predat,
  postdat,
  mod_post = mod_pre,
  cutoffs = NULL,
  SEM.pre = NULL,
  SEM.post = NULL,
  Fisher = FALSE,
  shiny = FALSE,
  main = "Test Scores",
  ...
)
```

Arguments

<code>mod_pre</code>	single-group model fitted by mirt . If not supplied the information will be extracted from the data input objects to compute the classical test theory version of the RCI statistics
<code>predat</code>	a vector (if one individual) or matrix/data.frame of response data to be scored, where each individuals' responses are included in exactly one row
<code>postdat</code>	same as <code>predat</code> , but with respect to the post/follow-up measurement
<code>mod_post</code>	(optional) IRT model for post-test if different from pre-test; otherwise, the pre-test model will be used
<code>cutoffs</code>	optional vector of length 2 indicating the type of cut-offs to report (e.g., <code>c(-1.96, 1.96)</code> reflects the 95 percent z-score type cut-off)
<code>SEM.pre</code>	standard error of measurement for the pretest. This can be used instead of <code>rxs.pre</code> and <code>SD.pre</code>
<code>SEM.post</code>	(optional) standard error of measurement for the post-test. Using this will create a pooled version of the SEM; otherwise, <code>SEM.post = SEM.pre</code>
<code>Fisher</code>	logical; use the Fisher/expected information function to compute the SE terms? If FALSE the SE information will be extracted from the select fscores method (default). Only applicable for unidimensional models

shiny	logical; launch an interactive shiny applications for real-time scoring of supplied total-scores or response vectors? Only requires mod_pre and (optional) mod_post inputs
main	main label to use when shiny=TRUE
...	additional arguments passed to fscores

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Jacobson, N. S., & Truax, P. (1991). Clinical significance: A statistical approach to defining meaningful change in psychotherapy research. *Journal of Consulting and Clinical Psychology*, 59, 12-19.
- Jabrayilov, R. , Emons, W. H. M., & Sijtsma, K. (2016). Comparison of Classical Test Theory and Item Response Theory in Individual Change Assessment. *Applied Psychological Measurement*, 40 (8), 559-572.

Examples

```
## Not run:

# simulate some data
N <- 1000
J <- 20      # number of items
a <- matrix(rlnorm(J,.2,.3))
d <- rnorm(J)

theta <- matrix(rnorm(N))
dat_pre <- simdata(a, d, itemtype = '2PL', Theta = theta)

# first 3 cases decrease by 1/2
theta2 <- theta - c(1/2, 1/2, 1/2, numeric(N-3))
dat_post <- simdata(a, d, itemtype = '2PL', Theta = theta2)

mod <- mirt(dat_pre)

# all changes using fitted model from pre data
RCI(mod, predat=dat_pre, postdat=dat_post)

# single response pattern change using EAP information
RCI(mod, predat=dat_pre[1,], postdat=dat_post[1,])

# WLE estimator with Fisher information for SE (see Jabrayilov et al. 2016)
RCI(mod, predat = dat_pre[1,], postdat = dat_post[1,],
    method = 'WLE', Fisher = TRUE)

# multiple respondents
```

```

RCI(mod, predat = dat_pre[1:6,], postdat = dat_post[1:6,])

# include large-sample z-type cutoffs
RCI(mod, predat = dat_pre[1:6,], postdat = dat_post[1:6,],
    cutoffs = c(-1.96, 1.96))

#####
# CTT version by omitting IRT model
# Requires either sample or population SEM's as input
(istats <- itemstats(dat_pre)$overall)
SEM.alpha <- istats$SEM.alpha # SEM estimate of dat_pre

# assumes SEM.post = SEM.pre
RCI(predat = dat_pre, postdat = dat_post, SEM.pre=SEM.alpha)

# include cutoffs
RCI(predat = dat_pre, postdat = dat_post, SEM.pre=SEM.alpha,
    cutoffs=c(-1.96, 1.96))

# allows SEM.post != SEM.pre
(istats.post <- itemstats(dat_post)$overall)
SEM.alpha.post <- istats.post$SEM.alpha

RCI(predat = dat_pre, postdat = dat_post,
    SEM.pre=SEM.alpha, SEM.post=SEM.alpha.post)

#####

# interactive shiny interfaces for live scoring
mod_pre <- mirt(Science)

# (optional) setup mod_post to have medium effect size change (d = 0.5)
sv <- mod2values(mod_pre)
sv$value[sv$name == 'MEAN_1'] <- 0.5
mod_post <- mirt(Science, pars=sv, TOL=NA)

# only use pre-test model for scoring
if(interactive()){
  RCI(mod_pre=mod_pre, shiny=TRUE)

  # use both pre-test and post-test models for including empirical priors
  RCI(mod_pre=mod_pre, mod_post=mod_post, shiny=TRUE,
      main='Perceptions of Science and Technology')
}

#####
# Example where individuals take completely different item set pre-post
# but prior calibration has been performed to equate the items

dat <- key2binary(SAT12,
    key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))

```

```

mod <- mirt(dat)

# with N=5 individuals under investigation
predat <- postdat <- dat[1:5,]
predat[, 17:32] <- NA
postdat[, 1:16] <- NA

head(predat)
head(postdat)

RCI(mod, predat, postdat)

## End(Not run)

```

read.mirt

Translate mirt parameters into suitable structure for plink package

Description

This function exports item parameters from the mirt package to the plink package.

Usage

```
read.mirt(x, as.irt.pars = TRUE, ...)
```

Arguments

x	a single object (or list of objects) returned from mirt, bfactor, or a single object returned by multipleGroup
as.irt.pars	if TRUE, the parameters will be output as an irt.pars object
...	additional arguments to be passed to coef()

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

Examples

```

## Not run:

## unidimensional
library(plink)

data <- expand.table(LSAT7)
(mod1 <- mirt(data, 1))
plinkpars <- read.mirt(mod1)
plot(plinkpars)
plot(mod1, type = 'trace')

```

```

# graded
mod2 <- mirt(Science, 1)
plinkpars <- read.mirt(mod2)
plot(plinkpars)
plot(mod2, type = 'trace')

# gpcm
mod3 <- mirt(Science, 1, itemtype = 'gpcm')
plinkpars <- read.mirt(mod3)
plot(plinkpars)
plot(mod3, type = 'trace')

# nominal
mod4 <- mirt(Science, 1, itemtype = 'nominal')
plinkpars <- read.mirt(mod4)
plot(plinkpars)
plot(mod4, type = 'trace')

## multidimensional

data <- expand.table(LSAT7)
(mod1 <- mirt(data, 2))
plinkpars <- read.mirt(mod1)
plinkpars
plot(plinkpars)
plot(mod1, type = 'trace')

cmmod <- mirt.model('
  F1 = 1,4,5
  F2 = 2-4')
model <- mirt(data, cmmod)
plot(read.mirt(model))
itemplot(model, 1)

# graded
mod2 <- mirt(Science, 2)
plinkpars <- read.mirt(mod2)
plinkpars
plot(plinkpars)
plot(mod2, type = 'trace')

### multiple group equating example
set.seed(1234)
dat <- expand.table(LSAT7)
group <- sample(c('g1', 'g2'), nrow(dat), TRUE)
dat1 <- dat[group == 'g1', ]
dat2 <- dat[group == 'g2', ]
mod1 <- mirt(dat1, 1)
mod2 <- mirt(dat2, 1)

# convert and combine pars
plinkMG <- read.mirt(list(g1=mod1, g2=mod2))

```

```
# equivalently:
# mod <- multipleGroup(dat, 1, group)
# plinkMG <- read.mirt(mod)

combine <- matrix(1:5, 5, 2)
comb <- combine.pars(plinkMG, combine, grp.names=unique(group))
out <- plink(comb, rescale="SL")
equate(out)
equate(out, method = 'OSE')

## End(Not run)
```

remap.distance

Remap item categories to have integer distances of 1

Description

The mirt package's estimation setup requires that all item responses have spaces equal to 1 (e.g., a Likert scale scored from 1 through 5). In the event that categories are missing the categories must be re-coded. This function is automatically called by the package estimation functions (e.g., [mirt](#)), however for convince this function has been extracted for users to better understand the remapping consequences.

Usage

```
remap.distance(data, message = TRUE)
```

Arguments

data	the response data to remap as a data.frame or matrix
message	logical; print message information pertaining to which items were remapped?

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```
# category 2 for item 1 missing
dat <- Science
dat[,1] <- ifelse(Science[,1] == 2, 1, Science[,1])
apply(dat, 2, table)

# mirt() automatically remaps categories
mod <- mirt(dat, 1)
coef(mod, simplify=TRUE)

# this is the transformed data used by mirt()
remap_dat <- remap.distance(dat)
apply(remap_dat, 2, table)
```

residuals-method

Compute model residuals

Description

Return model implied residuals for linear dependencies between items or at the person level. If the latent trait density was approximated (e.g., Davidian curves, Empirical histograms, etc) then passing `use_dentype_estimate = TRUE` will use the internally saved quadrature and density components (where applicable).

Usage

```
## S4 method for signature 'SingleGroupClass'
residuals(
  object,
  type = "LD",
  p.adjust = "none",
  df.p = FALSE,
  approx.z = FALSE,
  full.scores = FALSE,
  QMC = FALSE,
  printvalue = NULL,
  tables = FALSE,
  verbose = TRUE,
  Theta = NULL,
  suppress = NA,
  theta_lim = c(-6, 6),
  quadpts = NULL,
  fold = TRUE,
  upper = TRUE,
  technical = list(),
  ...
)
```

Arguments

object	an object of class <code>SingleGroupClass</code> or <code>MultipleGroupClass</code> . Bifactor models are automatically detected and utilized for better accuracy
type	type of residuals to be displayed. Can be either 'LD' or 'LDG2' for a local dependence matrix based on the X2 or G2 statistics (Chen & Thissen, 1997), 'Q3' for the statistic proposed by Yen (1984), 'JSI' for the jack-knife statistic proposed Edwards et al. (2018), 'exp' for the expected values for the frequencies of every response pattern, and 'expfull' for the expected values for every theoretically observable response pattern. For the 'LD' and 'LDG2' types, the upper diagonal elements represent the standardized residuals in the form of signed Cramers V coefficients
p.adjust	method to use for adjusting all p-values (see p.adjust for available options). Default is 'none'
df.p	logical; print the degrees of freedom and p-values?
approx.z	logical; transform $\chi^2(df)$ information from LD tests into approximate z-ratios instead using the transformation $z = \sqrt{2 * \chi^2} - \sqrt{2 * df - 1}$
full.scores	logical; compute relevant statistics for each subject in the original data?
QMC	logical; use quasi-Monte Carlo integration? If quadpts is omitted the default number of nodes is 5000
printvalue	a numeric value to be specified when using the <code>res='exp'</code> option. Only prints patterns that have standardized residuals greater than <code>abs(printvalue)</code> . The default (NULL) prints all response patterns
tables	logical; for LD type, return the observed, expected, and standardized residual tables for each item combination?
verbose	logical; allow information to be printed to the console?
Theta	a matrix of factor scores used for statistics that require empirical estimates (i.e., Q3). If supplied, arguments typically passed to <code>fscores()</code> will be ignored and these values will be used instead
suppress	a numeric value indicating which parameter local dependency combinations to flag as being too high (for LD, LDG2, and Q3 the standardize correlations are used; for JSI, the z-ratios are used). Absolute values for the standardized estimates greater than this value will be returned, while all values less than this value will be set to missing
theta_lim	range for the integration grid
quadpts	number of quadrature nodes to use. The default is extracted from model (if available) or generated automatically if not available
fold	logical; apply the sum 'folding' described by Edwards et al. (2018) for the JSI statistic?
upper	logical; which portion of the matrix (upper versus lower triangle) should the suppress argument be applied to?
technical	list of technical arguments when models are re-estimated (see mirt for details)
...	additional arguments to be passed to <code>fscores()</code>

References

- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Chen, W. H. & Thissen, D. (1997). Local dependence indices for item pairs using item response theory. *Journal of Educational and Behavioral Statistics*, 22, 265-289.
- Edwards, M. C., Houts, C. R. & Cai, L. (2018). A Diagnostic Procedure to Detect Departures From Local Independence in Item Response Theory Models. *Psychological Methods*, 23, 138-149.
- Yen, W. (1984). Effects of local item dependence on the fit and equating performance of the three parameter logistic model. *Applied Psychological Measurement*, 8, 125-145.

Examples

```
## Not run:

x <- mirt(Science, 1)
residuals(x)
residuals(x, tables = TRUE)
residuals(x, type = 'exp')
residuals(x, suppress = .15)
residuals(x, df.p = TRUE)
residuals(x, df.p = TRUE, p.adjust = 'fdr') # apply FWE control

# Pearson's X2 estimate for goodness-of-fit
full_table <- residuals(x, type = 'expfull')
head(full_table)
X2 <- with(full_table, sum((freq - exp)^2 / exp))
df <- nrow(full_table) - extract.mirt(x, 'nest') - 1
p <- pchisq(X2, df = df, lower.tail=FALSE)
data.frame(X2, df, p, row.names='Pearson-X2')

# above FOG test as a function
PearsonX2 <- function(x){
  full_table <- residuals(x, type = 'expfull')
  X2 <- with(full_table, sum((freq - exp)^2 / exp))
  df <- nrow(full_table) - extract.mirt(x, 'nest') - 1
  p <- pchisq(X2, df = df, lower.tail=FALSE)
  data.frame(X2, df, p, row.names='Pearson-X2')
}
PearsonX2(x)

# extract results manually
out <- residuals(x, df.p = TRUE, verbose=FALSE)
str(out)
out$df.p[1,2]

# with and without supplied factor scores
Theta <- fscores(x)
residuals(x, type = 'Q3', Theta=Theta)
residuals(x, type = 'Q3', method = 'ML')
```

```

# Edwards et al. (2018) JSI statistic
N <- 250
a <- rnorm(10, 1.7, 0.3)
d <- rnorm(10)
dat <- simdata(a, d, N=250, itemtype = '2PL')

mod <- mirt(dat, 1)
residuals(mod, type = 'JSI')
residuals(mod, type = 'JSI', fold=FALSE) # unfolded

# LD between items 1-2
aLD <- numeric(10)
aLD[1:2] <- rnorm(2, 2.55, 0.15)
a2 <- cbind(a, aLD)
dat <- simdata(a2, d, N=250, itemtype = '2PL')

mod <- mirt(dat, 1)

# JSI executed in parallel over multiple cores
if(interactive()) mirtCluster()
residuals(mod, type = 'JSI')

## End(Not run)

```

reverse.score

*Reverse score one or more items from a response matrix***Description**

Reverse score specific items given empirical range or specific scoring range.

Usage

```
reverse.score(data, which, range = NULL, append = ".RS")
```

Arguments

data	an object of class <code>data.frame</code> , <code>matrix</code> , or <code>table</code> with the response patterns
which	names of items or column integer location in data that should be rescored. If missing the all columns in data will be reverse scored
range	(optional) a named list to specify the low and high score ranges. Specified names must match the names found in data, and each element of this list should contain only two values. If items specified in which are omitted from this list then the empirical min/max information will be used instead
append	character vector indicating what to append to the item names that have been rescored

Value

returns the original data object with the specified items reverse scored replacing the original scoring scheme

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
a <- rlnorm(20)
a[c(1,5,10)] <- -a[c(1,5,10)]
diffs <- t(apply(matrix(runif(20*4, .3, 1), 20), 1, cumsum))
diffs <- -(diffs - rowMeans(diffs))
d <- diffs + rnorm(20)
dat <- simdata(a,d,itemtype='graded', N=300)
head(dat)

## Not run:
# fitted model has negative slopes due to flipped scoring
mod <- mirt(dat)
coef(mod, simplify=TRUE)$items
plot(mod, type = 'itemscore')

## End(Not run)

# reverse the scoring for items 1, 5, and 10 only using empirical min/max
revdat <- reverse.score(dat, c('Item_1', 'Item_5', 'Item_10'))
head(revdat)

# compare
apply(dat[,c(1,5,10)], 2, table)
apply(revdat[,c(1,5,10)], 2, table)

## Not run:
# slopes all positive now
mod2 <- mirt(revdat)
coef(mod2, simplify=TRUE)$items
plot(mod2, type = 'itemscore')

## End(Not run)

# use different empirical scoring information due to options not used
# 0 score not observed for item 1, though should have been rescored to a 4
dat[dat[,1] == 0, 1] <- 1
table(dat[,1])
```

```
# 4 score not observed for item 5, though should have been rescored to a 0
dat[dat[,5] == 4, 5] <- 3
table(dat[,5])

# specify theoretical scoring values in the range list
revdat2 <- reverse.score(dat, c('Item_1', 'Item_5', 'Item_10'),
                           range = list(Item_1 = c(0,4), Item_5 = c(0,4)))

head(revdat2)
table(dat[,1])
table(revdat2[,1])

table(dat[,5])
table(revdat2[,5])
```

RMSD_DIF

RMSD effect size statistic to quantify category-level DIF

Description

This function computes a set of RMSD "badness-of-fit" statistics when investigating DIF across a set of grouping variables. In a first step, a (potentially highly constrained) multiple group model is fitted, while in a second step the item (and person) parameters are estimated based on all examines across all groups. Category level DIF is assessed based on how well the pseudo-table of counts match the (constrained) probability functions implied by the original multiple group model (while also weighing across the implied density function of the latent traits). If the RMSD fit is poor, indicating non-ignorable DIF, then the multiple-group model should be adjusted to better account for the large response bias due to using a pooled model. See Lee and von Davier (2020) and Buchholz and Hartig (2019) for details.

Usage

```
RMSD_DIF(pooled_mod, flag = 0, probfun = TRUE, dentype = "norm")
```

Arguments

pooled_mod	a multiple-group model (used to compute the model-implied probability in the goodness-of-fit test)
flag	a numeric value used as a cut-off to help flag larger RMSD values (e.g., flag = .03 will highlight only categories with RMSD values greater than .03)
probfun	logical; use probability functions to compute RMSD? If FALSE, the expected score functions will be integrated instead, which may be useful for collapsing across the categories in polytomous items
dentype	density to use for the latent trait. Can be 'norm' to use a normal Gaussian density where the mean/variance are extracted from the model object(default), 'snorm' for a standard normal distribution, or 'empirical' to use the density estimate obtained via the E-table

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Buchholz, J., and Hartig, J. (2019). Comparing Attitudes Across Groups: An IRT-Based Item-Fit Statistic for the Analysis of Measurement Invariance. *Applied Psychological Measurement*, 43(3), 241-250. doi:[10.1177/0146621617748323](https://doi.org/10.1177/0146621617748323)

Lee, S. S., and von Davier, M. (2020). Improving measurement properties of the PISA home possessions scale through partial invariance modeling. *Psychological test and assessment modeling*, 62(1):55-83.

See Also

[DIF](#), [DRF](#), [multipleGroup](#), [empirical_ES](#)

Examples

```
## Not run:

#----- generate some data
set.seed(12345)
a <- a2 <- matrix(abs(rnorm(15,1,.3)), ncol=1)
d <- d2 <- matrix(rnorm(15,0,.7),ncol=1)

# item 1 has DIF
d2[1] <- d[1] - .5
a2[1] <- a[1] + 1

itemtype <- rep('2PL', nrow(a))
N <- 1000
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a2, d2, N, itemtype)
dat <- rbind(dataset1, dataset2)
group <- c(rep('D1', N), rep('D2', N))

#-----

# fully pooled model
pooled_mod <- multipleGroup(dat, 1, group=group,
  invariance = c(colnames(dat), 'free_mean', 'free_var'))
coef(pooled_mod, simplify=TRUE)

RMSD_DIF(pooled_mod)
RMSD_DIF(pooled_mod, dentype = 'empirical')
RMSD_DIF(pooled_mod, flag = .03)

# more freely estimated model (item 1 has 2 parameters estimated)
MGmod <- multipleGroup(dat, 1, group=group,
  invariance = c(colnames(dat)[-1], 'free_mean', 'free_var'))
coef(MGmod, simplify=TRUE)
```

```

# RMSD in item.1 now reduced (MG model accounts for DIF)
RMSD_DIF(MGmod)
RMSD_DIF(MGmod, flag = .03)

#####
# NA placeholders included when groups do not respond to specific items

a1 <- a2 <- rlnorm(20)
d <- d2 <- rnorm(20)
# item 5 contains DIF
a2[5] <- a1[5] + 1
d2[5] <- d[5] - 1/2
g <- rbeta(20, 5, 17)

dat1 <- simdata(a1, d, guess = g, N=1000, itemtype = '3PL')
dat1[, 11:13] <- NA # items 11:13 items NA for g1
dat2 <- simdata(a2, d2, guess = g, N=1000, itemtype = '3PL',
  mu=1/4, sigma=matrix(.75))
dat2[, 1:3] <- NA # items 1:3 items NA for g2
dat <- rbind(dat1, dat2)
group <- c(rep('g1', 1000), rep('g2', 1000))

mod <- multipleGroup(dat, "Theta = 1-20
  PRIOR = (1-20, g, norm, -1, 0.5)",
  group=group, itemtype='3PL',
  invariance = c(colnames(dat), 'free_mean', 'free_var'))
coef(mod, simplify = TRUE)

RMSD_DIF(mod)
RMSD_DIF(mod, flag = .03)

#####
# polytomous example
set.seed(12345)
a <- a2 <- matrix(rlnorm(20,.2,.3))

# for the graded model, ensure that there is enough space between the intercepts,
# otherwise closer categories will not be selected often (minimum distance of 0.3 here)
diffs <- t(apply(matrix(runif(20*4, .3, 1), 20), 1, cumsum))
diffs <- -(diffs - rowMeans(diffs))
d <- d2 <- diffs + rnorm(20)

# item 1 has slope + dif for first intercept parameter
d2[1] <- d[1] - .5
a2[1] <- a[1] + 1

itemtype <- rep('graded', nrow(a))
N <- 1000
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a2, d2, N, itemtype)
dat <- rbind(dataset1, dataset2)
group <- c(rep('D1', N), rep('D2', N))

```

```

#-----

# fully pooled model
pooled_mod <- multipleGroup(dat, 1, group=group,
                           invariance = c(colnames(dat), 'free_mean', 'free_var'))
coef(pooled_mod, simplify=TRUE)

# Item_1 fits poorly in several categories (RMSD > .05)
RMSD_DIF(pooled_mod)
RMSD_DIF(pooled_mod, flag = .05)
RMSD_DIF(pooled_mod, flag = .1, probfun = FALSE) # use expected score function

# more freely estimated model (item 1 has more parameters estimated)
MGmod <- multipleGroup(dat, 1, group=group,
                      invariance = c(colnames(dat)[-1], 'free_mean', 'free_var'))
coef(MGmod, simplify=TRUE)

# RMSDs in Item_1 now reduced (MG model better accounts for DIF)
RMSD_DIF(MGmod)
RMSD_DIF(MGmod, flag = .05)
RMSD_DIF(MGmod, probfun = FALSE, flag = .1) # use expected score function

## End(Not run)

```

SAT12

Description of SAT12 data

Description

Data obtained from the TESTFACT (Woods et al., 2003) manual, with 32 response pattern scored items for a grade 12 science assessment test (SAT) measuring topics of chemistry, biology, and physics. The scoring key for these data is [1, 4, 5, 2, 3, 1, 2, 1, 3, 1, 2, 4, 2, 1, 5, 3, 4, 4, 1, 4, 3, 3, 4, 1, 3, 5, 1, 3, 1, 5, 4, 5], respectively. However, careful analysis using the nominal response model suggests that the scoring key for item 32 may be incorrect, and should be changed from 5 to 3.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Wood, R., Wilson, D. T., Gibbons, R. D., Schilling, S. G., Muraki, E., & Bock, R. D. (2003). TESTFACT 4 for Windows: Test Scoring, Item Statistics, and Full-information Item Factor Analysis [Computer software]. Lincolnwood, IL: Scientific Software International.

Examples

```
## Not run:

itemstats(SAT12, use_ts = FALSE)

# score the data (missing scored as 0)
head(SAT12)
dat <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
head(dat)
itemstats(dat)

# score the data, missing (value of 8) treated as NA
SAT12missing <- SAT12
SAT12missing[SAT12missing == 8] <- NA
dat <- key2binary(SAT12missing,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,5))
head(dat)

# potentially better scoring for item 32 (based on nominal model finding)
dat <- key2binary(SAT12,
  key = c(1,4,5,2,3,1,2,1,3,1,2,4,2,1,5,3,4,4,1,4,3,3,4,1,3,5,1,3,1,5,4,3))

## End(Not run)
```

Science

*Description of Science data***Description**

A 4-item data set borrowed from ltm package in R, first example of the grm() function. See more complete documentation therein, as well as Karlheinz and Melich (1992).

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Karlheinz, R. and Melich, A. (1992). Euro-Barometer 38.1: *Consumer Protection and Perceptions of Science and Technology*. INRA (Europe), Brussels. [computer file]

Examples

```
## Not run:

itemstats(Science)

mod <- mirt(Science, 1)
plot(mod, type = 'trace')
```

```
## End(Not run)
```

secondOrderTest	<i>Second-order test of convergence</i>
-----------------	---

Description

Test whether terminated estimation criteria for a given model passes the second order test by checking the positive definiteness of the resulting Hessian matrix. This function, which accepts the symmetric Hessian/information matrix as the input, returns TRUE if the matrix is positive definite and FALSE otherwise.

Usage

```
secondOrderTest(mat, ..., method = "eigen")
```

Arguments

mat	symmetric matrix to test for positive definiteness (typically the Hessian at the highest point of model estimator, such as MLE or MAP)
...	arguments passed to either eigen , chol , or 'det' for the positiveness of the eigen values, positiveness of leading minors via the Cholesky decomposition, or evaluation of whether the determinant is greater than 0
method	method to use to test positive definiteness. Default is 'eigen'

Value

a matrix with all possible combinations

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```
## Not run:

# PD matrix
mod <- mirt(Science, 1, SE=TRUE)
info <- solve(vcov(mod)) ## observed information
secondOrderTest(info)
secondOrderTest(info, method = 'chol')
```

```
secondOrderTest(info, method = 'det')

# non-PD matrix
mat <- matrix(c(1,0,0,0,1,1,0,1,1), ncol=3)
mat
secondOrderTest(mat)
secondOrderTest(mat, method = 'chol')
secondOrderTest(mat, method = 'det')

## End(Not run)
```

show-method

Show model object

Description

Print model object summaries to the console.

Usage

```
## S4 method for signature 'SingleGroupClass'
show(object)
```

Arguments

object an object of class SingleGroupClass, MultipleGroupClass, or MixedClass

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```
## Not run:
x <- mirt(Science, 1)
show(x)

## End(Not run)
```

SIBTEST

*(Generalized) Simultaneous Item Bias Test (SIBTEST)***Description**

Classical test theory approach to detecting unidirectional and bidirectional (with one crossing location) DIF. This family of statistics is intended for unidimensional tests, and applies a regression-corrected matched-total score approach to quantify the response bias between two or more groups. Can be used for DIF, DBF, and DTF testing with two or more discrete groups.

Usage

```
SIBTEST(
  dat,
  group,
  suspect_set,
  match_set,
  focal_name = unique(group)[2],
  guess_correction = 0,
  Jmin = 5,
  na.rm = FALSE,
  randomize = FALSE,
  C = cbind(1, -diag(length(unique(group)) - 1L)),
  pairwise = FALSE,
  DIF = FALSE,
  p.adjust.method = "none",
  permute = 1000,
  pk_focal = FALSE,
  correction = TRUE,
  remove_cross = FALSE,
  details = FALSE,
  plot = "none",
  ...
)
```

Arguments

<code>dat</code>	integer-based dataset to be tested, containing dichotomous or polytomous responses
<code>group</code>	a (factor) vector indicating group membership with the same length as the number of rows in <code>dat</code>
<code>suspect_set</code>	an integer vector indicating which items to inspect with SIBTEST. Including only one value will perform a DIF test, while including more than one will perform a simultaneous bundle test (DBF); including all non-matched items will perform DTF. If missing, a simultaneous test using all the items not listed in <code>match_set</code> will be used (i.e., DTF)

<code>match_set</code>	an integer vector indicating which items to use as the items which are matched (i.e., contain no DIF). These are analogous to 'anchor' items in the likelihood method to locate DIF. If missing, all items other than the items found in the <code>suspect_set</code> will be used
<code>focal_name</code>	name of the focal group; e.g., 'focal'. If not specified then one will be selected automatically using <code>unique(group)[2]</code>
<code>guess_correction</code>	a vector of numbers from 0 to 1 indicating how much to correct the items for guessing. It's length should be the same as <code>ncol(dat)</code>
<code>Jmin</code>	the minimum number of observations required when splitting the data into focal and reference groups conditioned on the matched set
<code>na.rm</code>	logical; remove rows in <code>dat</code> with any missing values? If TRUE, rows with missing data will be removed, as well as the corresponding elements in the group input
<code>randomize</code>	logical; perform the crossing test for non-compensatory bias using Li and Stout's (1996) permutation approach? Default is FALSE, which uses the ad-hoc mixed degrees of freedom method suggested by Chalmers (2018)
<code>C</code>	a contrast matrix to use for pooled testing with more than two groups. Default uses an effects coding approach, where the last group (last column of the matrix) is treated as the reference group, and each column is associated with the respective name via <code>unique(group)</code> (i.e., the first column is the coefficient for <code>unique(group)[1]</code> , second column for <code>unique(group)[2]</code> , and so on)
<code>pairwise</code>	logical; perform pairwise comparisons in multi-group applications?
<code>DIF</code>	logical; should the elements in <code>suspect_set</code> be treated one at a time to test for DIF? Use of this logical will treat all other items as part of the <code>match_set</code> unless this input is provided explicitly. Default is FALSE to allow DBF and DTF tests
<code>p.adjust.method</code>	a character input dictating which method to use in <code>p.adjust</code> . when studying more than two groups. Default does not present any p-value adjustments
<code>permute</code>	number of permutations to perform when <code>randomize = TRUE</code> . Default is 1000
<code>pk_focal</code>	logical; using the group weights from the focal group instead of the total sample? Default is FALSE as per Shealy and Stout's recommendation
<code>correction</code>	logical; apply the composite correction for the difference between focal composite scores using the true-score regression technique? Default is TRUE, reflecting Shealy and Stout's linear extrapolation method
<code>remove_cross</code>	logical; remove the subtest information associated with the approximate crossing location? If TRUE this reflects the CSIBTEST definition of Li and Stout (1996); if FALSE, this reflects the version of CSIBTEST utilized by Chalmers (2018). Only applicable in two-group settings (in multi-group this is fixed to FALSE)
<code>details</code>	logical; return a data.frame containing the details required to compute SIBTEST?
<code>plot</code>	a character input indicating the type of plot to construct. Options are 'none' (default), 'observed' for the scaled focal subtest scores against the matched subtest scores, 'weights' for the proportion weights used (i.e., the proportion of observations at each matched score), 'difference' for the difference between the scaled focal subtest scores against the matched subtest scores, and

'wdifference' for the conditional differences multiplied by each respective weight. Note that the last plot reflects the components used in SIBTEST, and therefore the sum of these plotted observations will equal the beta coefficient for SIBTEST

... additional plotting arguments to be passed

Details

SIBTEST is similar to the Mantel-Haenszel approach for detecting DIF but uses a regression correction based on the KR-20/coefficient alpha reliability index to correct the observed differences when the latent trait distributions are not equal. Function supports the standard SIBTEST for dichotomous and polytomous data (compensatory) and supports crossing DIF testing (i.e., non-compensatory/non-uniform) using the asymptotic sampling distribution version of the Crossing-SIBTEST (CSIBTEST) statistic described by Chalmers (2018) and the permutation method described by Li and Stout (1996). This function also supports the multi-group generalizations (GSIBTEST and GCSIBTEST) proposed by Chalmers and Zheng (2023), where users may specify alternative contrast matrices to evaluate specific comparisons between groups as well as perform joint hypothesis tests.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

- Chalmers, R. P. (2018). Improving the Crossing-SIBTEST statistic for detecting non-uniform DIF. *Psychometrika*, 83, 2, 376-386.
- Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06
- Chalmers, R. P. & Zheng, G. (2023). Multi-group Generalizations of SIBTEST and Crossing-SIBTEST. *Applied Measurement in Education*, 36(2), 171-191, doi:10.1080/08957347.2023.2201703.
- Chang, H. H., Mazzeo, J. & Roussos, L. (1996). DIF for Polytomously Scored Items: An Adaptation of the SIBTEST Procedure. *Journal of Educational Measurement*, 33, 333-353.
- Li, H.-H. & Stout, W. (1996). A new procedure for detection of crossing DIF. *Psychometrika*, 61, 647-677.
- Shealy, R. & Stout, W. (1993). A model-based standardization approach that separates true bias/DIF from group ability differences and detect test bias/DTF as well as item bias/DIF. *Psychometrika*, 58, 159-194.

Examples

```
## Not run:

set.seed(1234)
n <- 30
N <- 500
a <- matrix(1, n)
d <- matrix(rnorm(n), n)
```

```

group <- c(rep('reference', N), rep('focal', N*2))

## -----
# groups completely equal
dat1 <- simdata(a, d, N, itemtype = 'dich')
dat2 <- simdata(a, d, N*2, itemtype = 'dich')
dat <- rbind(dat1, dat2)

# DIF (all other items as anchors)
SIBTEST(dat, group, suspect_set = 6)

# Some plots depicting the above tests
SIBTEST(dat, group, suspect_set = 6, plot = 'observed')
SIBTEST(dat, group, suspect_set = 6, plot = 'weights')
SIBTEST(dat, group, suspect_set = 6, plot = 'wdifference')

# Include CSIBTEST with randomization method
SIBTEST(dat, group, suspect_set = 6, randomize = TRUE)

# remove crossing-location (identical to Li and Stout 1996 definition of CSIBTEST)
SIBTEST(dat, group, suspect_set = 6, randomize = TRUE, remove_cross=TRUE)

# DIF (specific anchors)
SIBTEST(dat, group, match_set = 1:5, suspect_set = 6)
SIBTEST(dat, group, match_set = 1:5, suspect_set = 6, randomize=TRUE)

# DBF (all and specific anchors, respectively)
SIBTEST(dat, group, suspect_set = 11:30)
SIBTEST(dat, group, match_set = 1:5, suspect_set = 11:30)

# DTF
SIBTEST(dat, group, suspect_set = 11:30)
SIBTEST(dat, group, match_set = 1:10) #equivalent

# different hyper pars
dat1 <- simdata(a, d, N, itemtype = 'dich')
dat2 <- simdata(a, d, N*2, itemtype = 'dich', mu = .5, sigma = matrix(1.5))
dat <- rbind(dat1, dat2)
SIBTEST(dat, group, 6:30)
SIBTEST(dat, group, 11:30)

# DIF testing with anchors 1 through 5
SIBTEST(dat, group, 6, match_set = 1:5)
SIBTEST(dat, group, 7, match_set = 1:5)
SIBTEST(dat, group, 8, match_set = 1:5)

# DIF testing with all other items as anchors
SIBTEST(dat, group, 6)
SIBTEST(dat, group, 7)
SIBTEST(dat, group, 8)

## -----
## systematic differing slopes and intercepts (clear DTF)

```

```

dat1 <- simdata(a, d, N, itemtype = 'dich')
dat2 <- simdata(a + c(numeric(15), rnorm(n-15, 1, .25)), d + c(numeric(15), rnorm(n-15, 1, 1)),
  N*2, itemtype = 'dich')
dat <- rbind(dat1, dat2)
SIBTEST(dat, group, 6:30)
SIBTEST(dat, group, 11:30)

# Some plots depicting the above tests
SIBTEST(dat, group, suspect_set = 11:30, plot = 'observed')
SIBTEST(dat, group, suspect_set = 11:30, plot = 'weights')
SIBTEST(dat, group, suspect_set = 11:30, plot = 'wdifference')

# DIF testing using valid anchors
SIBTEST(dat, group, suspect_set = 6, match_set = 1:5)
SIBTEST(dat, group, suspect_set = 7, match_set = 1:5)
SIBTEST(dat, group, suspect_set = 30, match_set = 1:5)

# test DIF using specific match_set
SIBTEST(dat, group, suspect_set = 6:30, match_set = 1:5, DIF=TRUE)

# test DIF using all-other-as-anchors method (not typically recommended)
SIBTEST(dat, group, suspect_set = 1:30, DIF=TRUE)

# randomization method is fairly poor when smaller matched-set used
SIBTEST(dat, group, suspect_set = 30, match_set = 1:5, randomize=TRUE)
SIBTEST(dat, group, suspect_set = 30, randomize=TRUE)

## -----
# three group SIBTEST test
set.seed(1234)
n <- 30
N <- 1000
a <- matrix(1, n)
d <- matrix(rnorm(n), n)
group <- c(rep('group1', N), rep('group2', N), rep('group3', N))

# groups completely equal
dat1 <- simdata(a, d, N, itemtype = 'dich')
dat2 <- simdata(a, d, N, itemtype = 'dich')
dat3 <- simdata(a, d, N, itemtype = 'dich')
dat <- rbind(dat1, dat2, dat3)

# omnibus test using effects-coding contrast matrix (default)
SIBTEST(dat, group, suspect_set = 6)
SIBTEST(dat, group, suspect_set = 6, randomize=TRUE)

# explicit contrasts
SIBTEST(dat, group, suspect_set = 6, randomize=TRUE,
  C = matrix(c(1,-1,0), 1))

# test all items for DIF
SIBTEST(dat, group, suspect_set = 1:ncol(dat), DIF=TRUE)
SIBTEST(dat, group, suspect_set = 16:ncol(dat), DIF=TRUE,

```

```

        match_set = 1:15) # specific anchors

# post-hoc between two groups only
pick <- group %in% c('group1', 'group2')
SIBTEST(subset(dat, pick), group[pick], suspect_set = 1:ncol(dat), DIF=TRUE)

# post-hoc pairwise comparison for all groups
SIBTEST(dat, group, suspect_set = 1:ncol(dat), DIF=TRUE, pairwise = TRUE)

## systematic differing slopes and intercepts
dat2 <- simdata(a + c(numeric(15), .5,.5,.5,.5,.5, numeric(10)),
               d + c(numeric(15), 0,.6,.7,.8,.9, numeric(10)),
               N, itemtype = 'dich')
dat <- rbind(dat1, dat2, dat3)

SIBTEST(dat, group, suspect_set = 16)
SIBTEST(dat, group, suspect_set = 16, randomize=TRUE)

SIBTEST(dat, group, suspect_set = 19)
SIBTEST(dat, group, suspect_set = 19, randomize=TRUE)

SIBTEST(dat, group, suspect_set = c(16, 19), DIF=TRUE)
SIBTEST(dat, group, suspect_set = c(16, 19), DIF=TRUE, pairwise=TRUE)

## End(Not run)

```

simdata

Simulate response patterns

Description

Simulates response patterns for compensatory and noncompensatory MIRT models from multivariate normally distributed factor (θ) scores, or from a user input matrix of θ 's.

Usage

```

simdata(
  a,
  d,
  N,
  itemtype,
  sigma = NULL,
  mu = NULL,
  guess = 0,
  upper = 1,
  rho = NULL,
  nominal = NULL,

```

```

t = NULL,
Theta = NULL,
gpcm_mats = list(),
returnList = FALSE,
model = NULL,
equal.K = TRUE,
which.items = NULL,
mins = 0,
lca_cats = NULL,
prob.list = NULL
)

```

Arguments

a	a matrix/vector of slope parameters. If slopes are to be constrained to zero then use NA or simply set them equal to 0
d	a matrix/vector of intercepts. The matrix should have as many columns as the item with the largest number of categories, and filled empty locations with NA. When a vector is used the test is assumed to consist only of dichotomous items (because only one intercept per item is provided). When itemtype = 'lca' intercepts will not be used
N	sample size
itemtype	a character vector of length nrow(a) (or 1, if all the item types are the same) specifying the type of items to simulate. Inputs can either be the same as the inputs found in the itemtype argument in mirt or the internal classes defined by the package. Typical itemtype inputs that are passed to mirt are used then these will be converted into the respective internal classes automatically. If the internal class of the object is specified instead, the inputs can be 'dich', 'graded', 'gpcm', 'sequential', 'nominal', 'nestlogit', 'partcomp', 'gumm', 'lca', and all the models under the Luo (2001) family (see mirt), for dichotomous, graded, generalized partial credit, sequential, nominal, nested logit, partially compensatory, generalized graded unfolding model, latent class analysis model, and ordered unfolding models. Note that for the gpcm, nominal, and nested logit models there should be as many parameters as desired categories, however to parametrize them for meaningful interpretation the first category intercept should equal 0 for these models (second column for 'nestlogit', since first column is for the correct item traceline). For nested logit models the 'correct' category is always the lowest category (i.e., == 1). It may be helpful to use mod2values on data-sets that have already been estimated to understand the itemtypes more intimately
sigma	a covariance matrix of the underlying distribution. Default is the identity matrix. Used when Theta is not supplied
mu	a mean vector of the underlying distribution. Default is a vector of zeros. Used when Theta is not supplied
guess	a vector of guessing parameters for each item; only applicable for dichotomous items. Must be either a scalar value that will affect all of the dichotomous items, or a vector with as many values as to be simulated items

upper	same as guess, but for upper bound parameters
rho	a matrix of rho values to be used for the Lui (2001) family of ordered unfolding models (see mirt) to control the latitude of acceptance. All values must be positive
nominal	a matrix of specific item category slopes for nominal models. Should be the dimensions as the intercept specification with one less column, with NA in locations where not applicable. Note that during estimation the first slope will be constrained to 0 and the last will be constrained to the number of categories minus 1, so it is best to set these as the values for the first and last categories as well
t	matrix of t-values for the 'ggum' itemtype, where each row corresponds to a given item. Also determines the number of categories, where NA can be used for non-applicable categories
Theta	a user specified matrix of the underlying ability parameters, where <code>nrow(Theta) == N</code> and <code>ncol(Theta) == ncol(a)</code> . When this is supplied the N input is not required
gpcm_mats	a list of matrices specifying the scoring scheme for generalized partial credit models (see mirt for details)
returnList	logical; return a list containing the data, item objects defined by <code>mirt</code> containing the population parameters and item structure, and the latent trait matrix Theta? Default is FALSE
model	a single group object, typically returned by functions such as mirt or bfactor . Supplying this will render all other parameter elements (excluding the Theta, N, mu, and sigma inputs) redundant (unless explicitly provided). This input can therefore be used to create parametric bootstrap data whereby plausible data implied by the estimated model can be generated and evaluated
equal.K	logical; when a model input is supplied, should the generated data contain the same number of categories as the original data indicated by <code>extract.mirt(model, 'K')</code> ? Default is TRUE, which will redrawn data until this condition is satisfied
which.items	an integer vector used to indicate which items to simulate when a model input is included. Default simulates all items
mins	an integer vector (or single value to be used for each item) indicating what the lowest category should be. If model is supplied then this will be extracted from <code>slot(mod, 'Data')\$mins</code> , otherwise the default is 0
lca_cats	a vector indicating how many categories each lca item should have. If not supplied then it is assumed that 2 categories should be generated for each item
prob.list	an optional list containing matrix/data.frames of probabilities values for each category to be simulated. This is useful when creating customized probability functions to be sampled from

Details

Returns a data matrix simulated from the parameters, or a list containing the data, item objects, and Theta matrix.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Reckase, M. D. (2009). *Multidimensional Item Response Theory*. New York: Springer.

Examples

```
### Parameters from Reckase (2009), p. 153

set.seed(1234)

a <- matrix(c(
  .7471, .0250, .1428,
  .4595, .0097, .0692,
  .8613, .0067, .4040,
  1.0141, .0080, .0470,
  .5521, .0204, .1482,
  1.3547, .0064, .5362,
  1.3761, .0861, .4676,
  .8525, .0383, .2574,
  1.0113, .0055, .2024,
  .9212, .0119, .3044,
  .0026, .0119, .8036,
  .0008, .1905, 1.1945,
  .0575, .0853, .7077,
  .0182, .3307, 2.1414,
  .0256, .0478, .8551,
  .0246, .1496, .9348,
  .0262, .2872, 1.3561,
  .0038, .2229, .8993,
  .0039, .4720, .7318,
  .0068, .0949, .6416,
  .3073, .9704, .0031,
  .1819, .4980, .0020,
  .4115, 1.1136, .2008,
  .1536, 1.7251, .0345,
  .1530, .6688, .0020,
  .2890, 1.2419, .0220,
  .1341, 1.4882, .0050,
  .0524, .4754, .0012,
  .2139, .4612, .0063,
  .1761, 1.1200, .0870), 30, 3, byrow=TRUE)*1.702

d <- matrix(c(.1826, -.1924, -.4656, -.4336, -.4428, -.5845, -1.0403,
  .6431, .0122, .0912, .8082, -.1867, .4533, -1.8398, .4139,
  -.3004, -.1824, .5125, 1.1342, .0230, .6172, -.1955, -.3668,
  -1.7590, -.2434, .4925, -.3410, .2896, .006, .0329), ncol=1)*1.702
```

```

mu <- c(-.4, -.7, .1)
sigma <- matrix(c(1.21,.297,1.232,.297,.81,.252,1.232,.252,1.96),3,3)

dataset1 <- simdata(a, d, 2000, itemtype = '2PL')
dataset2 <- simdata(a, d, 2000, itemtype = '2PL', mu = mu, sigma = sigma)

#mod <- mirt(dataset1, 3, method = 'MHRM')
#coef(mod)

## Not run:

### Unidimensional graded response model with 5 categories each

a <- matrix(rlnorm(20,.2,.3))

# for the graded model, ensure that there is enough space between the intercepts,
# otherwise closer categories will not be selected often (minimum distance of 0.3 here)
diffs <- t(apply(matrix(runif(20*4, .3, 1), 20), 1, cumsum))
diffs <- -(diffs - rowMeans(diffs))
d <- diffs + rnorm(20)

dat <- simdata(a, d, 500, itemtype = 'graded')
# mod <- mirt(dat, 1)

### An example of a mixed item, bifactor loadings pattern with correlated specific factors

a <- matrix(c(
.8,.4,NA,
.4,.4,NA,
.7,.4,NA,
.8,NA,.4,
.4,NA,.4,
.7,NA,.4),ncol=3,byrow=TRUE)

d <- matrix(c(
-1.0,NA,NA,
1.5,NA,NA,
0.0,NA,NA,
0.0,-1.0,1.5, #the first 0 here is the recommended constraint for nominal
0.0,1.0,-1, #the first 0 here is the recommended constraint for gpcm
2.0,0.0,NA),ncol=3,byrow=TRUE)

nominal <- matrix(NA, nrow(d), ncol(d))
# the first 0 and last (ncat - 1) = 2 values are the recommended constraints
nominal[4, ] <- c(0,1.2,2)

sigma <- diag(3)
sigma[2,3] <- sigma[3,2] <- .25
items <- c('2PL','2PL','2PL','nominal','gpcm','graded')

dataset <- simdata(a,d,2000,items,sigma=sigma,nominal=nominal)

```

```

#mod <- bfactor(dataset, c(1,1,1,2,2,2), itemtype=c(rep('2PL', 3), 'nominal', 'gpcm','graded'))
#coef(mod)

#### Convert standardized factor loadings to slopes

F2a <- function(F, D=1.702){
  h2 <- rowSums(F^2)
  a <- (F / sqrt(1 - h2)) * D
  a
}

(F <- matrix(c(rep(.7, 5), rep(.5,5))))
(a <- F2a(F))
d <- rnorm(10)

dat <- simdata(a, d, 5000, itemtype = '2PL')
mod <- mirt(dat, 1)
coef(mod, simplify=TRUE)$items
summary(mod)

mod2 <- mirt(dat, 'F1 = 1-10
                CONSTRAIN = (1-5, a1), (6-10, a1)')
summary(mod2)
anova(mod2, mod)

#### Convert classical 3PL parameterization into slope-intercept form
nitems <- 50
as <- rlnorm(nitems, .2, .2)
bs <- rnorm(nitems, 0, 1)
gs <- rbeta(nitems, 5, 17)

# convert first item (only intercepts differ in resulting transformation)
traditional2mirt(c('a'=as[1], 'b'=bs[1], 'g'=gs[1], 'u'=1), cls='3PL')

# convert all difficulties to intercepts
ds <- numeric(nitems)
for(i in 1:nitems)
  ds[i] <- traditional2mirt(c('a'=as[i], 'b'=bs[i], 'g'=gs[i], 'u'=1),
                           cls='3PL')[2]

dat <- simdata(as, ds, N=5000, guess=gs, itemtype = '3PL')

# estimate with beta prior for guessing parameters
# mod <- mirt(dat, model="Theta = 1-50
#          PRIOR = (1-50, g, expbeta, 5, 17)", itemtype = '3PL')
# coef(mod, simplify=TRUE, IRTpars=TRUE)$items
# data.frame(as, bs, gs, us=1)

#### Unidimensional nonlinear factor pattern

theta <- rnorm(2000)
Theta <- cbind(theta, theta^2)

```

```

a <- matrix(c(
  .8,.4,
  .4,.4,
  .7,.4,
  .8,NA,
  .4,NA,
  .7,NA),ncol=2,byrow=TRUE)
d <- matrix(rnorm(6))
itemtype <- rep('2PL',6)

nonlindata <- simdata(a=a, d=d, itemtype=itemtype, Theta=Theta)

#model <- '
#F1 = 1-6
#(F1 * F1) = 1-3'
#mod <- mirt(nonlindata, model)
#coef(mod)

#### 2PLNRM model for item 4 (with 4 categories), 2PL otherwise

a <- matrix(rlnorm(4,0,.2))

# first column of item 4 is the intercept for the correct category of 2PL model,
# otherwise nominal model configuration
d <- matrix(c(
  -1.0,NA,NA,NA,
  1.5,NA,NA,NA,
  0.0,NA,NA,NA,
  1, 0.0,-0.5,0.5),ncol=4,byrow=TRUE)

nominal <- matrix(NA, nrow(d), ncol(d))
nominal[4, ] <- c(NA,0,.5,.6)

items <- c(rep('2PL',3),'nestlogit')

dataset <- simdata(a,d,2000,items,nominal=nominal)

#mod <- mirt(dataset, 1, itemtype = c('2PL', '2PL', '2PL', '2PLNRM'), key=c(NA,NA,NA,0))
#coef(mod)
#itemplot(mod,4)

# return list of simulation parameters
listobj <- simdata(a,d,2000,items,nominal=nominal, returnList=TRUE)
str(listobj)

# generate dataset from converged model
mod <- mirt(Science, 1, itemtype = c(rep('gpcm', 3), 'nominal'))
sim <- simdata(model=mod, N=1000)
head(sim)

Theta <- matrix(rnorm(100))
sim <- simdata(model=mod, Theta=Theta)

```

```

head(sim)

# alternatively, define a suitable object with functions from the mirtCAT package
# help(generate.mirt_object)
library(mirtCAT)

nitems <- 50
a1 <- rlnorm(nitems, .2, .2)
d <- rnorm(nitems)
g <- rbeta(nitems, 20, 80)
pars <- data.frame(a1=a1, d=d, g=g)
head(pars)

obj <- generate.mirt_object(pars, '3PL')
dat <- simdata(N=200, model=obj)

#### 10 item GGUMs test with 4 categories each
a <- rlnorm(10, .2, .2)
b <- rnorm(10) #passed to d= input, but used as the b parameters
diffs <- t(apply(matrix(runif(10*3, .3, 1), 10), 1, cumsum))
t <- -(diffs - rowMeans(diffs))

dat <- simdata(a, b, 1000, 'ggum', t=t)
apply(dat, 2, table)
# mod <- mirt(dat, 1, 'ggum')
# coef(mod)

### 10 items with the hyperbolic cosine model with differing category lengths
a <- matrix(1, 10)
d <- rnorm(10)
rho <- matrix(1:2, nrow=10, ncol=2, byrow=TRUE)
rho[1:2,2] <- NA # first two items have K=2 categories

dat <- simdata(a, d, 1000, 'hcm', rho=rho)
itemstats(dat)
# mod <- mirt(dat, 1, 'hcm')
# list(est=coef(mod, simplify=TRUE)$items, pop=cbind(a, d, log(rho)))

#####
# prob.list example

# custom probability function that returns a matrix
fun <- function(a, b, theta){
  P <- 1 / (1 + exp(-a * (theta-b)))
  cbind(1-P, P)
}

set.seed(1)
theta <- matrix(rnorm(100))
prob.list <- list()
nitems <- 5
a <- rlnorm(nitems, .2, .2); b <- rnorm(nitems, 0, 1/2)

```

```

for(i in 1:nitems) prob.list[[i]] <- fun(a[i], b[i], theta)
str(prob.list)

dat <- simdata(prob.list=prob.list)
head(dat)

# prob.list input is useful when defining custom items as well
name <- 'old2PL'
par <- c(a = .5, b = -2)
est <- c(TRUE, TRUE)
P.old2PL <- function(par,Theta, ncat){
  a <- par[1]
  b <- par[2]
  P1 <- 1 / (1 + exp(-1*a*(Theta - b)))
  cbind(1-P1, P1)
}

x <- createItem(name, par=par, est=est, P=P.old2PL)

prob.list[[1]] <- x@P(x@par, theta)

## End(Not run)

```

SingleGroupClass-class

Class "SingleGroupClass"

Description

Defines the object returned from `mirt` when model is exploratory.

Slots

Call: function call

Data: list of data, sometimes in different forms

Options: list of estimation options

Fit: a list of fit information

Model: a list of model-based information

ParObjects: a list of the S4 objects used during estimation

OptimInfo: a list of arguments from the optimization process

Internals: a list of internal arguments for secondary computations (inspecting this object is generally not required)

vcov: a matrix represented the asymptotic covariance matrix of the parameter estimates

time: a data.frame indicating the breakdown of computation times in seconds

Methods

```

anova signature(object = "SingleGroupClass")
coef signature(object = "SingleGroupClass")
plot signature(x = "SingleGroupClass", y = "missing")
print signature(x = "SingleGroupClass")
residuals signature(object = "SingleGroupClass")
show signature(object = "SingleGroupClass")
summary signature(object = "SingleGroupClass")

```

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

SLF

Social Life Feelings Data

Description

A 5-item data set analyzed by Bartholomew (1998). Data contains dichotomous responses (endorsement vs non-endorsement) from 1490 German respondents to five statements on perceptions of social life.

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Bartholomew, D., J. (1998). Scaling unobservable constructs in social science. *Journal of the Royal Statistical Society - Series C*, 47, 1-13.

Examples

```

## Not run:
# tabular format
data(SLF)
SLF

# full dataset
full <- expand.table(SLF)
itemstats(full)

```

```
mod <- mirt(full)
plot(mod, type = 'trace')

## End(Not run)
```

summary-method	<i>Summary of model object</i>
----------------	--------------------------------

Description

Transforms coefficients into a standardized factor loading’s metric. For MixedClass objects, the fixed and random coefficients are printed. Note that while the output to the console is rounded to three digits, the returned list of objects is not. For simulations, use `output <- summary(mod, verbose = FALSE)` to suppress the console messages.

Usage

```
## S4 method for signature 'SingleGroupClass'
summary(
  object,
  SE = TRUE,
  rotate = "oblimin",
  Target = NULL,
  suppress = 0,
  suppress.cor = 0,
  verbose = TRUE,
  ...
)
```

Arguments

object	an object of class SingleGroupClass, MultipleGroupClass, or MixedClass
SE	logical; include the standard errors for the standardized loadings? Requires the initial model to have included and estimated of the asymptotic covariance matrix (via, for instance, <code>mirt(..., SE = TRUE)</code>). If TRUE SEs are computed using the delta method
rotate	a string indicating which rotation to use for exploratory models, primarily from the GPArotation package (see documentation therein). Rotations currently supported are: 'promax', 'oblimin', 'varimax', 'quartimin', 'targetT', 'targetQ', 'pstT', 'pstQ', 'oblimax', 'entropy', 'quartimax', 'simplimax', 'bentlerT', 'bentlerQ', 'tandemI', 'tandemII', 'geominT', 'geominQ', 'cfT', 'cfQ', 'infomaxT', 'infomaxQ', 'mccammon', 'bifactorT', 'bifactorQ'. For models that are not exploratory this input will automatically be set to 'none'

Target	a dummy variable matrix indicting a target rotation pattern. This is required for rotations such as 'targetT', 'targetQ', 'pstT', and 'pstQ'
suppress	a numeric value indicating which (possibly rotated) factor loadings should be suppressed. Typical values are around .3 in most statistical software. Default is 0 for no suppression
suppress.cor	same as suppress, but for the correlation matrix output
verbose	logical; allow information to be printed to the console?
...	additional arguments to be passed

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

See Also

[coef-method](#)

Examples

```
## Not run:
x <- mirt(Science, 2)
summary(x)
summary(x, rotate = 'varimax')

## End(Not run)
```

testinfo

Function to calculate test information

Description

Given an estimated model compute the test information.

Usage

```
testinfo(
  x,
  Theta,
  degrees = NULL,
  group = NULL,
  individual = FALSE,
  which.items = 1:extract.mirt(x, "nitems")
)
```

Arguments

x	an object of class 'SingleGroupClass', or an object of class 'MultipleGroup-Class' if a suitable group input were supplied
Theta	a matrix of latent trait values
degrees	a vector of angles in degrees that are between 0 and 90. Only applicable when the input object is multidimensional
group	group argument to pass to <code>extract.group</code> function. Required when the input object is a multiple-group model
individual	logical; return a data.frame of information traceline for each item?
which.items	an integer vector indicating which items to include in the expected information function. Default uses all possible items

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
dat <- expand.table(deAyala)
(mirt(dat, 1, '2PL', pars = 'values'))
mod <- mirt(dat, 1, '2PL', constrain = list(c(1,5,9,13,17)))

Theta <- matrix(seq(-4,4,.01))
tinfo <- testinfo(mod, Theta)
plot(Theta, tinfo, type = 'l')

## Not run:

# compare information loss between two tests
tinfo_smaller <- testinfo(mod, Theta, which.items = 3:5)

# removed item informations
plot(Theta, iteminfo(extract.item(mod, 1), Theta), type = 'l')
plot(Theta, iteminfo(extract.item(mod, 2), Theta), type = 'l')

# most loss of info around -1 when removing items 1 and 2; expected given item info functions
plot(Theta, tinfo_smaller - tinfo, type = 'l')

## End(Not run)
```

`thetaComb`*Create all possible combinations of vector input*

Description

This function constructs all possible k-way combinations of an input vector. It is primarily useful when used in conjunction with the `mdirt` function, though users may have other uses for it as well. See [expand.grid](#) for more flexible combination formats.

Usage

```
thetaComb(theta, nfact, intercept = FALSE)
```

Arguments

<code>theta</code>	the vector from which all possible combinations should be obtained
<code>nfact</code>	the number of observations (and therefore the number of columns to return in the matrix of combinations)
<code>intercept</code>	logical; should a vector of 1's be appended to the first column of the result to include an intercept design component? Default is FALSE

Value

a matrix with all possible combinations

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
# all possible joint combinations for the vector -4 to 4
thetaComb(-4:4, 2)

# all possible binary combinations for four observations
thetaComb(c(0,1), 4)

# all possible binary combinations for four observations (with intercept)
thetaComb(c(0,1), 4, intercept=TRUE)
```

traditional2mirt	<i>Convert traditional IRT metric into slope-intercept form used in mirt</i>
------------------	--

Description

This is a helper function for users who have previously available traditional/classical IRT parameters and want to know the equivalent slope-intercept translation used in `mirt`. Note that this function assumes that the supplied models are unidimensional by definition (i.e., will have only one slope/discrimination) and in the logistic metric (i.e., logistic-ogive scaling coefficient $D=1$). If there is no supported slope-intercept transformation available then the original vector of parameters will be returned by default.

Usage

```
traditional2mirt(x, cls, ncat)
```

Arguments

<code>x</code>	a vector of parameters to transform
<code>cls</code>	the class or itemtype of the supplied model
<code>ncat</code>	the number of categories implied by the IRT model

Details

Supported class transformations for the `cls` input are:

Rasch, 2PL, 3PL, 3PLu, 4PL Form must be: (discrimination, difficulty, lower-bound, upper-bound)

graded Form must be: (discrimination, difficulty 1, difficulty 2, ..., difficulty $k-1$)

gpcm Form must be: (discrimination, difficulty 1, difficulty 2, ..., difficulty $k-1$)

nominal Form must be: (discrimination 1, discrimination 2, ..., discrimination k , difficulty 1, difficulty 2, ..., difficulty k)

Value

a named vector of slope-intercept parameters (if supported)

Examples

```
# classical 3PL model
vec <- c(a=1.5, b=-1, g=.1, u=1)
slopeint <- traditional2mirt(vec, '3PL', ncat=2)
slopeint

# classical graded model (four category)
vec <- c(a=1.5, b1=-1, b2=0, b3=1.5)
slopeint <- traditional2mirt(vec, 'graded', ncat=4)
slopeint
```

```
# classical generalize partial credit model (four category)
vec <- c(a=1.5, b1=-1, b2=0, b3=1.5)
slopeint <- traditional2mirt(vec, 'gpcm', ncat=4)
slopeint

# classical nominal model (4 category)
vec <- c(a1=.5, a2 = -1, a3=1, a4=-.5, d1=1, d2=-1, d3=-.5, d4=.5)
slopeint <- traditional2mirt(vec, 'nominal', ncat=4)
slopeint
```

vcov-method

Extract parameter variance covariance matrix

Description

Extract parameter variance covariance matrix

Usage

```
## S4 method for signature 'SingleGroupClass'
vcov(object)
```

Arguments

object an object of class SingleGroupClass, MultipleGroupClass, or MixedClass

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:[10.18637/jss.v048.i06](https://doi.org/10.18637/jss.v048.i06)

Examples

```
## Not run:
x <- mirt(Science, 1, SE=TRUE)
vcov(x)

## End(Not run)
```

wald

*Wald statistics for mirt models***Description**

Compute a Wald test given an L vector or matrix of numeric contrasts. Requires that the model information matrix be computed (by passing SE = TRUE when estimating the model). Use `wald(model)` to observe how the information matrix columns are named, especially if the estimated model contains constrained parameters (e.g., IPL).

Usage

```
wald(object, L, C = NULL)
```

Arguments

object	estimated object from <code>mirt</code> , <code>bfactor</code> , <code>multipleGroup</code> , <code>mixedmirt</code> , or <code>mdirt</code>
L	a coefficient matrix with dimensions <code>nconstrasts x npars.estimated</code> , or a character vector giving the hypothesis in symbolic form (syntax format borrowed from the <code>car</code> package; see Details below). Omitting this value will return the column names of the information matrix used to identify the (potentially constrained) parameters
C	a constant vector of population parameters to be compared along side L, where <code>length(C) == row(L)</code>. By default a vector of 0's is constructed. Note that when using the syntax input for L this argument is ignored The following description is borrowed from <code>car</code> package documentation pertaining to the character vector input to the argument L: "The hypothesis matrix can be supplied as a numeric matrix (or vector), the rows of which specify linear combinations of the model coefficients, which are tested equal to the corresponding entries in the right-hand-side vector, which defaults to a vector of zeroes. Alternatively, the hypothesis can be specified symbolically as a character vector with one or more elements, each of which gives either a linear combination of coefficients, or a linear equation in the coefficients (i.e., with both a left and right side separated by an equals sign). Components of a linear expression or linear equation can consist of numeric constants, or numeric constants multiplying coefficient names (in which case the number precedes the coefficient, and may be separated from it by spaces or an asterisk); constants of 1 or -1 may be omitted. Spaces are always optional. Components are separated by plus or minus signs. Newlines or tabs in hypotheses will be treated as spaces. See the examples below."

Author(s)

Phil Chalmers <rphilip.chalmers@gmail.com>

References

Chalmers, R., P. (2012). mirt: A Multidimensional Item Response Theory Package for the R Environment. *Journal of Statistical Software*, 48(6), 1-29. doi:10.18637/jss.v048.i06

Examples

```
## Not run:

# View parnumber index
data(LSAT7)
data <- expand.table(LSAT7)
mod <- mirt(data, 1, SE = TRUE)
coef(mod)

# see how the information matrix relates to estimated parameters, and how it lines up
# with the parameter index
(infonames <- wald(mod))
index <- mod2values(mod)
index[index$est, ]

# second item slope equal to 0?
L <- matrix(0, 1, 10)
L[1,3] <- 1
wald(mod, L)

# same as above using character syntax input
infonames
wald(mod, "a1.5 = 0")

# simultaneously test equal factor slopes for item 1 and 2, and 4 and 5
L <- matrix(0, 2, 10)
L[1,1] <- L[2, 7] <- 1
L[1,3] <- L[2, 9] <- -1
L
wald(mod, L)

# Again, using more efficient syntax
infonames
wald(mod, c("a1.1 = a1.5", "a1.13 = a1.17"))

# log-Likelihood tests (requires estimating a new model)
cmodel <- 'theta = 1-5
          CONSTRAIN = (1,2, a1), (4,5, a1)'
mod2 <- mirt(data, cmodel)
# or, equivalently
#mod2 <- mirt(data, 1, constrain = list(c(1,5), c(13,17)))
anova(mod2, mod)

#####
# test equality of means in multi-group model:
# H0: (mu1 - mu2) = (mu3 - mu4)
```

```
set.seed(12345)
a <- matrix(abs(rnorm(15,1,.3)), ncol=1)
d <- matrix(rnorm(15,0,.7),ncol=1)
itemtype <- rep('2PL', nrow(a))
N <- 500
dataset1 <- simdata(a, d, N, itemtype)
dataset2 <- simdata(a, d, N, itemtype, mu = .5)
dataset3 <- simdata(a, d, N, itemtype, mu = -1)
dataset4 <- simdata(a, d, N, itemtype, mu = -.5)
dat <- rbind(dataset1, dataset2, dataset3, dataset4)
group <- factor(rep(paste0('D', 1:4), each=N))
levels(group)
models <- 'F1 = 1-15'

# 3 means estimated
mod_free <- multipleGroup(dat, models, group = group, SE=TRUE,
                          invariance=c('slopes', 'intercepts', 'free_var','free_means'))
wald(mod_free) # obtain parameter names
# View(mod2values(mod_free))

# reference group mean = 0 by default
wald(mod_free, c("0 - MEAN_1.123 = MEAN_1.185 - MEAN_1.247"))

## End(Not run)
```

Index

- * **Lagrange**
 - lagrange, 95
- * **SIBTEST**
 - SIBTEST, 204
- * **area**
 - areainfo, 7
- * **bootstrapped**
 - boot.mirt, 20
- * **bootstrap**
 - boot.LR, 18
- * **change**
 - RCI, 186
- * **classes**
 - DiscreteClass-class, 36
 - MixedClass-class, 149
 - MixtureClass-class, 158
 - MultipleGroupClass-class, 171
 - SingleGroupClass-class, 217
- * **conversion**
 - likert2int, 97
- * **convert**
 - mod2values, 159
- * **createGroup**
 - createGroup, 23
- * **createItem**
 - createItem, 25
- * **data**
 - ASVAB, 8
 - Attitude, 9
 - Bock1997, 18
 - deAyala, 29
 - expand.table, 56
 - imputeMissing, 77
 - itemstats, 92
 - likert2int, 97
 - LSAT6, 100
 - LSAT7, 101
 - poly2dich, 181
 - SAT12, 200
 - Science, 201
 - simdata, 209
 - SLF, 218
- * **derivatives**
 - numerical_deriv, 172
- * **differential**
 - DIF, 30
 - DRF, 38
 - DTF, 45
- * **discrimination**
 - MDIFF, 105
 - MDISC, 113
- * **effects**
 - fixef, 68
 - randef, 185
- * **empirical**
 - empirical_plot, 51
- * **errors**
 - boot.mirt, 20
- * **expected**
 - expected.item, 58
 - expected.test, 59
- * **extract**
 - extract.group, 61
 - extract.item, 62
 - extract.mirt, 63
- * **factor.scores**
 - fscores, 70
- * **fit**
 - itemfit, 79
 - itemGAM, 84
 - M2, 102
 - personfit, 173
- * **fixed**
 - fixef, 68
- * **functioning**
 - DIF, 30
 - DRF, 38
 - DTF, 45

- * **imputation**
 - averageMI, 10
 - * **impute**
 - imputeMissing, 77
 - * **index**
 - RCI, 186
 - * **information**
 - areainfo, 7
 - iteminfo, 88
 - testinfo, 220
 - * **item**
 - DIF, 30
 - itemfit, 79
 - itemGAM, 84
 - * **likelihood**
 - PLCI.mirt, 175
 - * **models**
 - bfactor, 12
 - mdirt, 106
 - mirt, 114
 - multipleGroup, 160
 - * **model**
 - M2, 102
 - mod2values, 159
 - * **multiple**
 - averageMI, 10
 - * **numerical**
 - numerical_deriv, 172
 - * **package**
 - mirt-package, 4
 - * **parallel**
 - mirtCluster, 148
 - * **parametric**
 - boot.LR, 18
 - * **person**
 - personfit, 173
 - * **plink**
 - read.mirt, 189
 - * **plots**
 - empirical_plot, 51
 - * **plot**
 - itemplot, 89
 - * **profiled**
 - PLCI.mirt, 175
 - * **random**
 - randef, 185
 - * **reliability**
 - empirical_rxx, 53
 - marginal_rxx, 104
 - * **reliable**
 - RCI, 186
 - * **response**
 - DRF, 38
 - * **scores**
 - estfun.AllModelClass, 54
 - * **score**
 - expected.test, 59
 - * **standard**
 - boot.mirt, 20
 - * **test**
 - DTF, 45
 - lagrange, 95
 - * **tracelines**
 - probtrace, 184
 - * **value**
 - expected.item, 58
 - * **wald**
 - wald, 225
- anova, 13
- anova, DiscreteClass-method
(anova-method), 5
- anova, MixedClass-method (anova-method),
5
- anova, MixtureClass-method
(anova-method), 5
- anova, MultipleGroupClass-method
(anova-method), 5
- anova, SingleGroupClass-method
(anova-method), 5
- anova-method, 5, 125
- areainfo, 7
- ASVAB, 8
- Attitude, 9
- averageMI, 10, 74, 131
- bfactor, 12, 54, 55, 62, 90, 123, 131, 144,
146, 184, 211
- Bock1997, 18
- boot, 37
- boot.LR, 18
- boot.mirt, 20, 109, 125, 153, 176
- bs, 116, 121
- chol, 202
- coef, DiscreteClass-method
(coef-method), 21

- coef, MixedClass-method (coef-method), 21
- coef, MixtureClass-method (coef-method), 21
- coef, MultipleGroupClass-method (coef-method), 21
- coef, SingleGroupClass-method (coef-method), 21
- coef-method, 21, 125
- createGroup, 23
- createItem, 25, 102, 106, 108, 116, 125
- data.matrix, 115
- deAyala, 29
- DeltaMethod, 29
- DIF, 30, 41, 45, 46, 160, 162, 198
- DiscreteClass-class, 36
- draw_parameters, 37, 39
- DRF, 33, 37, 38, 162, 198
- DTF, 45
- eigen, 202
- empirical_ES, 48, 198
- empirical_plot, 51, 93
- empirical_rxx, 53, 104
- estfun.AllModelClass, 54
- expand.grid, 222
- expand.table, 56, 131
- expected.item, 58, 60
- expected.test, 59, 59
- extract.group, 7, 61, 62, 65, 104–106, 113, 114, 162, 184, 221
- extract.item, 59, 61, 62, 65, 88, 131, 184
- extract.mirt, 61, 62, 63, 122, 131, 159
- factor, 161
- fixedCalib, 65
- fixef, 68, 131, 153
- fscores, 49, 53, 54, 70, 78, 80, 81, 85, 102, 108, 109, 124, 125, 173, 179, 186, 187
- gen.difficulty, 76
- imputeMissing, 77, 125
- integrate, 7
- itemfit, 79, 86, 108, 109, 125, 174
- itemGAM, 52, 79, 82, 84
- iteminfo, 88, 90, 131
- itemplot, 52, 89, 125
- itemstats, 52, 92, 131
- key2binary, 94, 97, 131
- lagrange, 95
- lattice, 40, 49, 52, 81, 85, 86, 91, 179
- likert2int, 97
- logLik, DiscreteClass-method (logLik-method), 99
- logLik, MixedClass-method (logLik-method), 99
- logLik, MixtureClass-method (logLik-method), 99
- logLik, MultipleGroupClass-method (logLik-method), 99
- logLik, SingleGroupClass-method (logLik-method), 99
- logLik-method, 99
- LSAT6, 100
- LSAT7, 101
- M2, 102, 108, 109, 125
- marginal_rxx, 54, 104
- MDIFF, 105
- mdirt, 36, 54, 55, 71, 73, 106, 180, 222
- MDISC, 106, 113
- mirt, 12–14, 20, 25, 52, 54, 55, 65, 66, 69, 71, 72, 76, 96, 107, 108, 114, 146, 151–153, 161, 162, 184, 186, 191, 193, 210, 211, 217
- mirt-package, 4
- mirt.model, 12, 13, 66, 107, 109, 115, 119, 124, 125, 143, 151, 161
- mirtCluster, 18, 31, 73, 109, 118, 122, 125, 148, 150
- MixedClass-class, 149, 153
- mixedmirt, 69, 123, 124, 131, 146, 149, 150, 185
- MixtureClass-class, 158
- mod2values, 65, 96, 131, 159, 176, 210
- multipleGroup, 13, 31, 33, 41, 46, 49, 54, 55, 61, 64–66, 123, 131, 146, 158, 160, 171, 198
- MultipleGroupClass-class, 13, 162, 171
- ns, 116, 121
- numerical_deriv, 26, 29, 96, 172
- optim, 118

- p.adjust, [32](#), [40](#), [80](#), [193](#), [205](#)
- personfit, [82](#), [125](#), [173](#)
- PLCI.mirt, [175](#)
- plogis, [146](#)
- plot, DiscreteClass, missing-method
 - (plot, MultipleGroupClass, missing-method), [177](#)
- plot, MixtureClass, missing-method
 - (plot, MultipleGroupClass, missing-method), [177](#)
- plot, MultipleGroupClass, missing-method, [177](#)
- plot, MultipleGroupClass-method
 - (plot, MultipleGroupClass, missing-method), [177](#)
- plot, SingleGroupClass, missing-method
 - (plot, MultipleGroupClass, missing-method), [177](#)
- plot, SingleGroupClass-method
 - (plot, MultipleGroupClass, missing-method), [177](#)
- plot-method, [125](#)
- plot-method
 - (plot, MultipleGroupClass, missing-method), [177](#)
- plot.itemGAM (itemGAM), [84](#)
- poly2dich, [97](#), [181](#)
- print, DiscreteClass-method
 - (print-method), [182](#)
- print, MixedClass-method (print-method), [182](#)
- print, MixtureClass-method
 - (print-method), [182](#)
- print, MultipleGroupClass-method
 - (print-method), [182](#)
- print, SingleGroupClass-method
 - (print-method), [182](#)
- print-method, [182](#)
- print.mirt_df, [182](#)
- print.mirt_list, [183](#)
- print.mirt_matrix, [183](#)
- probtrace, [131](#), [184](#)
- randef, [153](#), [185](#)
- RCI, [186](#)
- read.mirt, [189](#)
- remap.distance, [191](#)
- residuals, DiscreteClass-method
 - (residuals-method), [192](#)
- residuals, MixtureClass-method
 - (residuals-method), [192](#)
- residuals, MultipleGroupClass-method
 - (residuals-method), [192](#)
- residuals, SingleGroupClass-method
 - (residuals-method), [192](#)
- residuals-method, [125](#), [192](#)
- reverse.score, [195](#)
- RMSD_DIF, [197](#)
- SAT12, [200](#)
- Science, [201](#)
- secondOrderTest, [202](#)
- show, DiscreteClass-method
 - (show-method), [203](#)
- show, MixedClass-method (show-method), [203](#)
- show, MixtureClass-method (show-method), [203](#)
- show, MultipleGroupClass-method
 - (show-method), [203](#)
- show, SingleGroupClass-method
 - (show-method), [203](#)
- show-method, [203](#)
- SIBTEST, [204](#)
- simdata, [131](#), [209](#)
- SingleGroupClass-class, [13](#), [123](#), [217](#)
- SLF, [218](#)
- summary, DiscreteClass-method
 - (summary-method), [219](#)
- summary, MixedClass-method
 - (summary-method), [219](#)
- summary, MixtureClass-method
 - (summary-method), [219](#)
- summary, MultipleGroupClass-method
 - (summary-method), [219](#)
- summary, SingleGroupClass-method
 - (summary-method), [219](#)
- summary-method, [125](#), [219](#)
- testinfo, [104](#), [131](#), [220](#)
- thetaComb, [107](#), [109](#), [222](#)
- traditional2mirt, [223](#)
- uniroot, [77](#)
- vcov, DiscreteClass-method
 - (vcov-method), [224](#)
- vcov, MixedClass-method (vcov-method), [224](#)

`vcov,MixtureClass-method` (`vcov-method`),
 [224](#)
`vcov,MultipleGroupClass-method`
 (`vcov-method`), [224](#)
`vcov,SingleGroupClass-method`
 (`vcov-method`), [224](#)
`vcov-method`, [224](#)

`wald`, [96](#), [109](#), [117](#), [125](#), [225](#)

`xyplot`, [49](#)