

Package ‘hBayesDM’

September 9, 2026

Title Hierarchical Bayesian Modeling of Decision-Making Tasks

Version 2.0.1

Date 2026-09-09

Description Fit an array of decision-making tasks with computational models in a hierarchical Bayesian framework. Can perform hierarchical Bayesian analysis of various computational models with a single line of coding (Ahn et al., 2017) <doi:10.1162/CPSY_a_00002>.

Depends R (>= 4.4), methods

Imports posterior (>= 1.6), bayesplot (>= 1.11), loo (>= 2.7), grid, parallel, ggplot2, data.table

Additional_repositories <https://stan-dev.r-universe.dev>

URL <https://github.com/CCS-Lab/hBayesDM>

BugReports <https://github.com/CCS-Lab/hBayesDM/issues>

License GPL-3

NeedsCompilation no

Encoding UTF-8

Collate 'preprocess_funcs.R' 'fit_cmdstan.R' 'settings.R'
'hBayesDM_model.R' 'alt_delta.R' 'alt_gamma.R'
'bandit2arm_delta.R' 'bandit4arm2_kalman_filter.R'
'bandit4arm_2par_lapse.R' 'bandit4arm_4par.R'
'bandit4arm_lapse.R' 'bandit4arm_lapse_decay.R'
'bandit4arm_singleA_lapse.R' 'banditNarm_2par_lapse.R'
'banditNarm_4par.R' 'banditNarm_delta.R'
'banditNarm_kalman_filter.R' 'banditNarm_lapse.R'
'banditNarm_lapse_decay.R' 'banditNarm_singleA_lapse.R'
'bart_ewmv.R' 'bart_par4.R' 'cgt_cm.R' 'choiceRT_ddm.R'
'choiceRT_ddm_single.R' 'choiceRT_lba.R'
'choiceRT_lba_single.R' 'cra_exp.R' 'cra_linear.R'
'dbdm_prob_weight.R' 'dd_cs.R' 'dd_cs_single.R' 'dd_exp.R'
'dd_hyperbolic.R' 'dd_hyperbolic_single.R' 'estimate_mode.R'
'extract_ic.R' 'gng_m1.R' 'gng_m2.R' 'gng_m3.R' 'gng_m4.R'

'hBayesDM.R' 'hdi.R' 'hgf_ibrb.R' 'hgf_ibrb_single.R'
 'igt_orl.R' 'igt_pvl_decay.R' 'igt_pvl_delta.R' 'igt_vpp.R'
 'multiplot.R' 'peer_ocu.R' 'plot.hBayesDM.R' 'plot_dist.R'
 'plot_hdi.R' 'plot_ind.R' 'print_fit.R' 'prl_ewa.R'
 'prl_fictitious.R' 'prl_fictitious_multipleB.R'
 'prl_fictitious_rp.R' 'prl_fictitious_rp_woa.R'
 'prl_fictitious_woa.R' 'prl_rp.R' 'prl_rp_multipleB.R'
 'pstRT_ddm.R' 'pstRT_rlddm1.R' 'pstRT_rlddm6.R' 'pst_Q.R'
 'pst_gainloss_Q.R' 'ra_noLA.R' 'ra_noRA.R' 'ra_prospect.R'
 'rdt_happiness.R' 'rhat.R' 'task2AFC_sdt.R' 'ts_par4.R'
 'ts_par6.R' 'ts_par7.R' 'ug_bayes.R' 'ug_delta.R' 'wcs_sql.R'
 'zzz.R'

Suggests cmdstanr (>= 0.8.1), testthat

Config/roxygen2/version 8.1.0

Author Woo-Young Ahn [aut, cre],
 Nate Haines [aut],
 Lei Zhang [aut],
 Jinwoo Jeong [ctb],
 Harhim Park [ctb],
 Jaeyeong Yang [ctb],
 Jethro Lee [ctb]

Maintainer Woo-Young Ahn <wooyoung.ahn@gmail.com>

Repository CRAN

Date/Publication 2026-09-09 13:10:02 UTC

Contents

alt_delta	4
alt_gamma	7
bandit2arm_delta	11
bandit4arm2_kalman_filter	14
bandit4arm_2par_lapse	18
bandit4arm_4par	21
bandit4arm_lapse	25
bandit4arm_lapse_decay	28
bandit4arm_singleA_lapse	32
banditNarm_2par_lapse	35
banditNarm_4par	39
banditNarm_delta	42
banditNarm_kalman_filter	46
banditNarm_lapse	49
banditNarm_lapse_decay	53
banditNarm_singleA_lapse	57
bart_ewmv	60
bart_par4	64
cgt_cm	67

choiceRT_ddm	71
choiceRT_ddm_single	75
cra_exp	78
cra_linear	82
dbdm_prob_weight	85
dd_cs	89
dd_cs_single	93
dd_exp	96
dd_hyperbolic	100
dd_hyperbolic_single	103
estimate_mode	107
extract_ic	108
gng_m1	108
gng_m2	112
gng_m3	115
gng_m4	119
hdi	123
hgf_ibrb	123
hgf_ibrb_single	127
igt_orl	131
igt_pvl_decay	135
igt_pvl_delta	138
igt_vpp	142
multiplot	146
peer_ocu	146
plot_dist	150
plot_hdi	151
plot_ind	152
print_fit	153
prl_ewa	154
prl_fictitious	157
prl_fictitious_multipleB	161
prl_fictitious_rp	164
prl_fictitious_rp_woa	168
prl_fictitious_woa	172
prl_rp	175
prl_rp_multipleB	179
pstRT_ddm	182
pstRT_rlddm1	186
pstRT_rlddm6	190
pst_gainloss_Q	194
pst_Q	197
ra_noLA	201
ra_noRA	205
ra_prospect	208
rdt_happiness	212
rhat	215
task2AFC_sdt	216

ts_par4 219

ts_par6 223

ts_par7 227

ug_bayes 231

ug_delta 234

wcs_sql 238

Index 242

alt_delta	<i>Rescorla-Wagner (Delta) Model</i>
-----------	--------------------------------------

Description

Hierarchical Bayesian Modeling of the Aversive Learning Task using Rescorla-Wagner (Delta) Model. It has the following parameters: A (learning rate), beta (inverse temperature), gamma (risk preference).

- **Task:** Aversive Learning Task (Browning et al., 2015)
- **Model:** Rescorla-Wagner (Delta) Model

Usage

```
alt_delta(  
  data = NULL,  
  niter = 4000,  
  nwarmup = 1000,  
  nchain = 4,  
  ncore = 1,  
  nthin = 1,  
  inits = "vb",  
  ind_pars = "mean",  
  model_regressor = FALSE,  
  vb = FALSE,  
  inc_postpred = FALSE,  
  adapt_delta = 0.95,  
  stepsize = 1,  
  max_treedepth = 10,  
  seed = 42,  
  ...  
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "sub-jID", "choice", "outcome", "bluePunish", "orangePunish". See Details below for more information.
------	---

niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Aversive Learning Task, there should be 5 columns of data with the labels "subjID", "choice", "outcome", "bluePunish", "orangePunish". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial (blue == 1, orange == 2).

outcome Integer value representing the outcome of the given trial (punishment == 1, and non-punishment == 0).

bluePunish Floating point value representing the magnitude of punishment for blue on that trial (e.g., 10, 97)

orangePunish Floating point value representing the magnitude of punishment for orange on that trial (e.g., 23, 45)

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Lili Zhang](#) <lili.zhang27@mail.dcu.ie>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("alt_delta").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Browning, M., Behrens, T. E., Jocham, G., O'reilly, J. X., & Bishop, S. J. (2015). Anxious individuals have difficulty learning the causal statistics of aversive environments. *Nature neuroscience*, 18(4), 590.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- alt_delta(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- alt_delta(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

alt_gamma

Rescorla-Wagner (Gamma) Model

Description

Hierarchical Bayesian Modeling of the Aversive Learning Task using Rescorla-Wagner (Gamma) Model. It has the following parameters: A (learning rate), beta (inverse temperature), gamma (risk preference).

- **Task:** Aversive Learning Task (Browning et al., 2015)
- **Model:** Rescorla-Wagner (Gamma) Model

Usage

```
alt_gamma(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "sub-ID", "choice", "outcome", "bluePunish", "orangePunish". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Aversive Learning Task, there should be 5 columns of data with the labels "subjID", "choice", "outcome", "bluePunish", "orangePunish". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial (blue == 1, orange == 2).

outcome Integer value representing the outcome of the given trial (punishment == 1, and non-punishment == 0).

bluePunish Floating point value representing the magnitude of punishment for blue on that trial (e.g., 10, 97)

orangePunish Floating point value representing the magnitude of punishment for orange on that trial (e.g., 23, 45)

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple

chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: `hBayesDM` 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Lili Zhang](#) <lili.zhang27@mail.dcu.ie>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("alt_gamma").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A `CmdStanMCMC` object (or `CmdStanVB` when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Browning, M., Behrens, T. E., Jocham, G., O'reilly, J. X., & Bishop, S. J. (2015). Anxious individuals have difficulty learning the causal statistics of aversive environments. *Nature neuroscience*, 18(4), 590.

See Also

We refer users to our in-depth tutorial for an example of using `hBayesDM`: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- alt_gamma(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)
```

```

# Run the model with example data
output <- alt_gamma(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

bandit2arm_delta	<i>Rescorla-Wagner (Delta) Model</i>
------------------	--------------------------------------

Description

Hierarchical Bayesian Modeling of the 2-Armed Bandit Task using Rescorla-Wagner (Delta) Model. It has the following parameters: A (learning rate), tau (inverse temperature).

- **Task:** 2-Armed Bandit Task (Erev et al., 2010; Hertwig et al., 2004)
- **Model:** Rescorla-Wagner (Delta) Model

Usage

```

bandit2arm_delta(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_tredepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the 2-Armed Bandit Task, there should be 3 columns of data with the labels "subjID", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1 or 2.

outcome Integer value representing the outcome of the given trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("bandit2arm_delta").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

- Erev, I., Ert, E., Roth, A. E., Haruvy, E., Herzog, S. M., Hau, R., et al. (2010). A choice prediction competition: Choices from experience and from description. *Journal of Behavioral Decision Making*, 23(1), 15-47. <https://doi.org/10.1002/bdm.683>
- Hertwig, R., Barron, G., Weber, E. U., & Erev, I. (2004). Decisions From Experience and the Effect of Rare Events in Risky Choice. *Psychological Science*, 15(8), 534-539. <https://doi.org/10.1111/j.0956-7976.2004.00715.x>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- bandit2arm_delta(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- bandit2arm_delta(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

bandit4arm2_kalman_filter

Kalman Filter

Description

Hierarchical Bayesian Modeling of the 4-Armed Bandit Task (modified) using Kalman Filter. It has the following parameters: `lambda` (decay factor), `theta` (decay center), `beta` (inverse softmax temperature), `mu0` (anticipated initial mean of all 4 options), `s0` (anticipated initial sd (uncertainty factor) of all 4 options), `sD` (sd of diffusion noise).

- **Task:** 4-Armed Bandit Task (modified)
- **Model:** Kalman Filter (Daw et al., 2006)

Usage

```
bandit4arm2_kalman_filter(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.

<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the 4-Armed Bandit Task (modified), there should be 3 columns of data with the labels "subjID", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, or 4.

outcome Integer value representing the outcome of the given trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Yoonseo Zoh](#) <<zohyos7@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("bandit4arm2_kalman_filter").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Daw, N. D., O'Doherty, J. P., Dayan, P., Seymour, B., & Dolan, R. J. (2006). Cortical substrates for exploratory decisions in humans. *Nature*, 441(7095), 876-879.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- bandit4arm2_kalman_filter(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- bandit4arm2_kalman_filter(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")
```

```

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

bandit4arm_2par_lapse *3 Parameter Model, without C (choice perseveration), R (reward sensitivity), and P (punishment sensitivity). But with xi (noise)*

Description

Hierarchical Bayesian Modeling of the 4-Armed Bandit Task using 3 Parameter Model, without C (choice perseveration), R (reward sensitivity), and P (punishment sensitivity). But with xi (noise). It has the following parameters: Arew (reward learning rate), Apun (punishment learning rate), xi (noise).

- **Task:** 4-Armed Bandit Task
- **Model:** 3 Parameter Model, without C (choice perseveration), R (reward sensitivity), and P (punishment sensitivity). But with xi (noise) (Aylward et al., 2018)

Usage

```

bandit4arm_2par_lapse(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_tredepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the 4-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, or 4.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("bandit4arm_2par_lapse").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Aylward, Valton, Ahn, Bond, Dayan, Roiser, & Robinson (2018) Altered decision-making under uncertainty in unmedicated mood and anxiety disorders. PsyArxiv. 10.31234/osf.io/k5b8m

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- bandit4arm_2par_lapse(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- bandit4arm_2par_lapse(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

bandit4arm_4par

4 Parameter Model, without C (choice perseveration)

Description

Hierarchical Bayesian Modeling of the 4-Armed Bandit Task using 4 Parameter Model, without C (choice perseveration). It has the following parameters: Arew (reward learning rate), Apun (punishment learning rate), R (reward sensitivity), P (punishment sensitivity).

- **Task:** 4-Armed Bandit Task
- **Model:** 4 Parameter Model, without C (choice perseveration) (Seymour et al., 2012)

Usage

```
bandit4arm_4par(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the 4-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, or 4.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("bandit4arm_4par").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Seymour, Daw, Roiser, Dayan, & Dolan (2012). Serotonin Selectively Modulates Reward Value in Human Decision-Making. *J Neuro*, 32(17), 5833-5842.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- bandit4arm_4par(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- bandit4arm_4par(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
```

```

rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

bandit4arm_lapse	<i>5 Parameter Model, without C (choice perseveration) but with xi (noise)</i>
------------------	--

Description

Hierarchical Bayesian Modeling of the 4-Armed Bandit Task using 5 Parameter Model, without C (choice perseveration) but with xi (noise). It has the following parameters: Arew (reward learning rate), Apun (punishment learning rate), R (reward sensitivity), P (punishment sensitivity), xi (noise).

- **Task:** 4-Armed Bandit Task
- **Model:** 5 Parameter Model, without C (choice perseveration) but with xi (noise) (Seymour et al., 2012)

Usage

```

bandit4arm_lapse(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_tredepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the 4-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, or 4.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("bandit4arm_lapse").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Seymour, Daw, Roiser, Dayan, & Dolan (2012). Serotonin Selectively Modulates Reward Value in Human Decision-Making. *J Neuro*, 32(17), 5833-5842.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- bandit4arm_lapse(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- bandit4arm_lapse(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

bandit4arm_lapse_decay

5 Parameter Model, without C (choice perseveration) but with xi (noise). Added decay rate (Niv et al., 2015, J. Neuro).

Description

Hierarchical Bayesian Modeling of the 4-Armed Bandit Task using 5 Parameter Model, without C (choice perseveration) but with xi (noise). Added decay rate (Niv et al., 2015, J. Neuro).. It has the following parameters: Arew (reward learning rate), Apun (punishment learning rate), R (reward sensitivity), P (punishment sensitivity), xi (noise), d (decay rate).

- **Task:** 4-Armed Bandit Task
- **Model:** 5 Parameter Model, without C (choice perseveration) but with xi (noise). Added decay rate (Niv et al., 2015, J. Neuro). (Aylward et al., 2018)

Usage

```
bandit4arm_lapse_decay(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the 4-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, or 4.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("bandit4arm_lapse_decay").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Aylward, Valton, Ahn, Bond, Dayan, Roiser, & Robinson (2018) Altered decision-making under uncertainty in unmedicated mood and anxiety disorders. PsyArxiv. 10.31234/osf.io/k5b8m

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- bandit4arm_lapse_decay(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- bandit4arm_lapse_decay(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
```

```

rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

bandit4arm_singleA_lapse

4 Parameter Model, without C (choice perseveration) but with xi (noise). Single learning rate both for R and P.

Description

Hierarchical Bayesian Modeling of the 4-Armed Bandit Task using 4 Parameter Model, without C (choice perseveration) but with xi (noise). Single learning rate both for R and P. It has the following parameters: A (learning rate), R (reward sensitivity), P (punishment sensitivity), xi (noise).

- **Task:** 4-Armed Bandit Task
- **Model:** 4 Parameter Model, without C (choice perseveration) but with xi (noise). Single learning rate both for R and P. (Aylward et al., 2018)

Usage

```

bandit4arm_singleA_lapse(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_tredepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the 4-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, or 4.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: **hBayesDM** 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("bandit4arm_singleA_lapse").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Aylward, Valton, Ahn, Bond, Dayan, Roiser, & Robinson (2018) Altered decision-making under uncertainty in unmedicated mood and anxiety disorders. PsyArxiv. 10.31234/osf.io/k5b8m

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- bandit4arm_singleA_lapse(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- bandit4arm_singleA_lapse(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

banditNarm_2par_lapse *3 Parameter Model, without C (choice perseveration), R (reward sensitivity), and P (punishment sensitivity). But with xi (noise)*

Description

Hierarchical Bayesian Modeling of the N-Armed Bandit Task using 3 Parameter Model, without C (choice perseveration), R (reward sensitivity), and P (punishment sensitivity). But with xi (noise). It has the following parameters: Arew (reward learning rate), Apun (punishment learning rate), xi (noise).

- **Task:** N-Armed Bandit Task
- **Model:** 3 Parameter Model, without C (choice perseveration), R (reward sensitivity), and P (punishment sensitivity). But with xi (noise) (Aylward et al., 2018)

Usage

```
banditNarm_2par_lapse(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_tredepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, it's possible to set model-specific argument(s) as follows: Narm Number of arms used in Multi-armed Bandit Task If not given, the number of unique choice will be used.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the N-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, ... N.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Cheol Jun Cho](#) <<cjfwndns1@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("banditNarm_2par_lapse").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Aylward, Valton, Ahn, Bond, Dayan, Roiser, & Robinson (2018) Altered decision-making under uncertainty in unmedicated mood and anxiety disorders. PsyArxiv. 10.31234/osf.io/k5b8m

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- banditNarm_2par_lapse(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- banditNarm_2par_lapse(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")
```

```

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

banditNarm_4par	<i>4 Parameter Model, without C (choice perseveration)</i>
-----------------	--

Description

Hierarchical Bayesian Modeling of the N-Armed Bandit Task using 4 Parameter Model, without C (choice perseveration). It has the following parameters: Arew (reward learning rate), Apun (punishment learning rate), R (reward sensitivity), P (punishment sensitivity).

- **Task:** N-Armed Bandit Task
- **Model:** 4 Parameter Model, without C (choice perseveration) (Seymour et al., 2012)

Usage

```

banditNarm_4par(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: Narm Number of arms used in Multi-armed Bandit Task If not given, the number of unique choice will be used.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the N-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, ... N.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Cheol Jun Cho](#) <<cjfwndns1@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("banditNarm_4par").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Seymour, Daw, Roiser, Dayan, & Dolan (2012). Serotonin Selectively Modulates Reward Value in Human Decision-Making. *J Neuro*, 32(17), 5833-5842.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- banditNarm_4par(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- banditNarm_4par(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

banditNarm_delta

Rescorla-Wagner (Delta) Model

Description

Hierarchical Bayesian Modeling of the N-Armed Bandit Task using Rescorla-Wagner (Delta) Model. It has the following parameters: A (learning rate), tau (inverse temperature).

- **Task:** N-Armed Bandit Task (Erev et al., 2010; Hertwig et al., 2004)
- **Model:** Rescorla-Wagner (Delta) Model

Usage

```
banditNarm_delta(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_tredepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, it's possible to set model-specific argument(s) as follows: Narm Number of arms used in Multi-armed Bandit Task If not given, the number of unique choice will be used.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the N-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, ... N.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Cheol Jun Cho](#) <<cjfwndns1@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("banditNarm_delta").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Erev, I., Ert, E., Roth, A. E., Haruvy, E., Herzog, S. M., Hau, R., et al. (2010). A choice prediction competition: Choices from experience and from description. *Journal of Behavioral Decision Making*, 23(1), 15-47. <https://doi.org/10.1002/bdm.683>

Hertwig, R., Barron, G., Weber, E. U., & Erev, I. (2004). Decisions From Experience and the Effect of Rare Events in Risky Choice. *Psychological Science*, 15(8), 534-539. <https://doi.org/10.1111/j.0956-7976.2004.00715.x>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- banditNarm_delta(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- banditNarm_delta(
```

```

data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

banditNarm_kalman_filter

Kalman Filter

Description

Hierarchical Bayesian Modeling of the N-Armed Bandit Task (modified) using Kalman Filter. It has the following parameters: lambda (decay factor), theta (decay center), beta (inverse softmax temperature), mu0 (anticipated initial mean of all 4 options), s0 (anticipated initial sd (uncertainty factor) of all 4 options), sD (sd of diffusion noise).

- **Task:** N-Armed Bandit Task (modified)
- **Model:** Kalman Filter (Daw et al., 2006)

Usage

```

banditNarm_kalman_filter(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == nthin$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: Narm Number of arms used in Multi-armed Bandit Task If not given, the number of unique choice will be used.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the N-Armed Bandit Task (modified), there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, ... N.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Yoonseo Zoh <zohyos7@gmail.com>](mailto:zohyos7@gmail.com), [Cheol Jun Cho <cjfwndns1@gmail.com>](mailto:cjfwndns1@gmail.com)

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("banditNarm_kalman_filter").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Daw, N. D., O’Doherty, J. P., Dayan, P., Seymour, B., & Dolan, R. J. (2006). Cortical substrates for exploratory decisions in humans. *Nature*, 441(7095), 876-879.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- banditNarm_kalman_filter(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- banditNarm_kalman_filter(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

banditNarm_lapse	<i>5 Parameter Model, without C (choice perseveration) but with xi (noise)</i>
------------------	--

Description

Hierarchical Bayesian Modeling of the N-Armed Bandit Task using 5 Parameter Model, without C (choice perseveration) but with xi (noise). It has the following parameters: Arew (reward learning rate), Apun (punishment learning rate), R (reward sensitivity), P (punishment sensitivity), xi (noise).

- **Task:** N-Armed Bandit Task
- **Model:** 5 Parameter Model, without C (choice perseveration) but with xi (noise) (Seymour et al., 2012)

Usage

```
banditNarm_lapse(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.

vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: Narm Number of arms used in Multi-armed Bandit Task If not given, the number of unique choice will be used.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the N-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, ... N.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple

chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: `hBayesDM` 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Cheol Jun Cho](#) <<cjfwndns1@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("banditNarm_lapse").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A `CmdStanMCMC` object (or `CmdStanVB` when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Seymour, Daw, Roiser, Dayan, & Dolan (2012). Serotonin Selectively Modulates Reward Value in Human Decision-Making. *J Neuro*, 32(17), 5833-5842.

See Also

We refer users to our in-depth tutorial for an example of using `hBayesDM`: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- banditNarm_lapse(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)
```

```
# Run the model with example data
output <- banditNarm_lapse(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

banditNarm_lapse_decay

5 Parameter Model, without C (choice perseveration) but with xi (noise). Added decay rate (Niv et al., 2015, J. Neuro).

Description

Hierarchical Bayesian Modeling of the N-Armed Bandit Task using 5 Parameter Model, without C (choice perseveration) but with xi (noise). Added decay rate (Niv et al., 2015, J. Neuro).. It has the following parameters: Arew (reward learning rate), Apun (punishment learning rate), R (reward sensitivity), P (punishment sensitivity), xi (noise), d (decay rate).

- **Task:** N-Armed Bandit Task
- **Model:** 5 Parameter Model, without C (choice perseveration) but with xi (noise). Added decay rate (Niv et al., 2015, J. Neuro). (Aylward et al., 2018)

Usage

```
banditNarm_lapse_decay(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
```

```

    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: Narm Number of arms used in Multi-armed Bandit Task If not given, the number of unique choice will be used.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the N-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, ... N.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Cheol Jun Cho](#) <<cjfwndns1@gmail.com>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("banditNarm_lapse_decay").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Aylward, Valton, Ahn, Bond, Dayan, Roiser, & Robinson (2018) Altered decision-making under uncertainty in unmedicated mood and anxiety disorders. PsyArxiv. 10.31234/osf.io/k5b8m

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- banditNarm_lapse_decay(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- banditNarm_lapse_decay(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

banditNarm_singleA_lapse

4 Parameter Model, without C (choice perseveration) but with xi (noise). Single learning rate both for R and P.

Description

Hierarchical Bayesian Modeling of the N-Armed Bandit Task using 4 Parameter Model, without C (choice perseveration) but with xi (noise). Single learning rate both for R and P. It has the following parameters: A (learning rate), R (reward sensitivity), P (punishment sensitivity), xi (noise).

- **Task:** N-Armed Bandit Task
- **Model:** 4 Parameter Model, without C (choice perseveration) but with xi (noise). Single learning rate both for R and P. (Aylward et al., 2018)

Usage

```
banditNarm_singleA_lapse(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.

nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: Narm Number of arms used in Multi-armed Bandit Task If not given, the number of unique choice will be used.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the N-Armed Bandit Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on the given trial: 1, 2, 3, ... N.

gain Floating point value representing the amount of currency won on the given trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on the given trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the

necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Cheol Jun Cho](#) <<cjfwndns1@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("banditNarm_singleA_lapse").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Aylward, Valton, Ahn, Bond, Dayan, Roiser, & Robinson (2018) Altered decision-making under uncertainty in unmedicated mood and anxiety disorders. PsyArxiv. 10.31234/osf.io/k5b8m

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- banditNarm_singleA_lapse(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- banditNarm_singleA_lapse(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

bart_ewmv

Exponential-Weight Mean-Variance Model

Description

Hierarchical Bayesian Modeling of the Balloon Analogue Risk Task using Exponential-Weight Mean-Variance Model. It has the following parameters: ϕ (prior belief of burst), η (updating exponent), ρ (risk preference), τ (inverse temperature), λ (loss aversion).

- **Task:** Balloon Analogue Risk Task
- **Model:** Exponential-Weight Mean-Variance Model (Park et al., 2020)

Usage

```
bart_ewmv(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
```

```

ncore = 1,
nthin = 1,
inits = "vb",
ind_pars = "mean",
model_regressor = FALSE,
vb = FALSE,
inc_postpred = FALSE,
adapt_delta = 0.95,
stepsize = 1,
max_treedepth = 10,
seed = 42,
...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "pumps", "explosion". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.

seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Balloon Analogue Risk Task, there should be 3 columns of data with the labels "subjID", "pumps", "explosion". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

pumps The number of pumps.

explosion 0: intact, 1: burst

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Harhim Park <<hrpark12@gmail.com>>, Jaeyeong Yang <<jaeyeong.yang1125@gmail.com>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("bart_ewmv").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Park, H., Yang, J., Vassileva, J., & Ahn, W. (2020). The Exponential-Weight Mean-Variance Model: A novel computational model for the Balloon Analogue Risk Task. <https://doi.org/10.31234/osf.io/sdzj4>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- bart_ewmv(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- bart_ewmv(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)
```

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

bart_par4

Re-parameterized version of BART model with 4 parameters

Description

Hierarchical Bayesian Modeling of the Balloon Analogue Risk Task using Re-parameterized version of BART model with 4 parameters. It has the following parameters: `phi` (prior belief of balloon not bursting), `eta` (updating rate), `gam` (risk-taking parameter), `tau` (inverse temperature).

- **Task:** Balloon Analogue Risk Task
- **Model:** Re-parameterized version of BART model with 4 parameters (van Ravenzwaaij et al., 2011)

Usage

```
bart_par4(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "pumps", "explosion". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.

ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Balloon Analogue Risk Task, there should be 3 columns of data with the labels "subjID", "pumps", "explosion". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

pumps The number of pumps.

explosion 0: intact, 1: burst

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: `hBayesDM` 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Harhim Park](#) <<hrpark12@gmail.com>>, [Jaeyeong Yang](#) <<jaeyeong.yang1125@gmail.com>>, [Ayoung Lee](#) <<aylee2008@naver.com>>, [Jeongbin Oh](#) <<ows0104@gmail.com>>, [Jiyeon Lee](#) <<nicole.lee2001@gmail.com>>, [Junha Jang](#) <<andy627robo@naver.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("bart_par4").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A `CmdStanMCMC` object (or `CmdStanVB` when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

van Ravenzwaaij, D., Dutilh, G., & Wagenmakers, E. J. (2011). Cognitive model decomposition of the BART: Assessment and application. *Journal of Mathematical Psychology*, 55(1), 94-105.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- bart_par4(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- bart_par4(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

cgt_cm

Cumulative Model

Description

Hierarchical Bayesian Modeling of the Cambridge Gambling Task using Cumulative Model. It has the following parameters: alpha (probability distortion), c (color bias), rho (relative loss sensitivity), beta (discounting rate), gamma (choice sensitivity).

- **Task:** Cambridge Gambling Task (Rogers et al., 1999)
- **Model:** Cumulative Model

Usage

```
cgt_cm(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
```

```

ncore = 1,
nthin = 1,
inits = "vb",
ind_pars = "mean",
model_regressor = FALSE,
vb = FALSE,
inc_postpred = FALSE,
adapt_delta = 0.95,
stepsize = 1,
max_treedepth = 10,
seed = 42,
...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "gamble_type", "percentage_staked", "trial_initial_points", "assessment_stage", "red_chosen", "n_red_boxes". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "y_hat_col", "y_hat_bet", "bet_utils".
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. Not available for this model.
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.

seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Cambridge Gambling Task, there should be 7 columns of data with the labels "subjID", "gamble_type", "percentage_staked", "trial_initial_points", "assessment_stage", "red_chosen", "n_red_boxes". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

gamble_type Integer value representing whether the bets on the current trial were presented in descending (0) or ascending (1) order.

percentage_staked Integer value representing the bet percentage (not proportion) selected on the current trial: 5, 25, 50, 75, or 95.

trial_initial_points Floating point value representing the number of points that the subject has at the start of the current trial (e.g., 100, 150, etc.).

assessment_stage Integer value representing whether the current trial is a practice trial (0) or a test trial (1). Only test trials are used for model fitting.

red_chosen Integer value representing whether the red color was chosen (1) versus the blue color (0).

n_red_boxes Integer value representing the number of red boxes shown on the current trial: 1, 2, 3, ..., or 9.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Nathaniel Haines](#) <<haines.175@osu.edu>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("cgt_cm").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Rogers, R. D., Everitt, B. J., Baldacchino, A., Blackshaw, A. J., Swainson, R., Wynne, K., Baker, N. B., Hunter, J., Carthy, T., London, M., Deakin, J. F. W., Sahakian, B. J., Robbins, T. W. (1999). Dissociable deficits in the decision-making cognition of chronic amphetamine abusers, opiate abusers, patients with focal damage to prefrontal cortex, and tryptophan-depleted normal volunteers: evidence for monoaminergic mechanisms. *Neuropsychopharmacology*, 20, 322–339.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- cgt_cm(
```

```

data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- cgt_cm(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

choiceRT_ddm

Drift Diffusion Model

Description

Hierarchical Bayesian Modeling of the Choice Reaction Time Task using Drift Diffusion Model. It has the following parameters: alpha (boundary separation), beta (bias), delta (drift rate), tau (non-decision time).

- **Task:** Choice Reaction Time Task
- **Model:** Drift Diffusion Model (Ratcliff, 1978)

Usage

```

choiceRT_ddm(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,

```

```

    max_treedepth = 10,
    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "RT". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. Not available for this model.
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: RTbound Floating point value representing the lower bound (i.e., minimum allowed) reaction time. Defaults to 0.1 (100 milliseconds).

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Choice Reaction Time Task, there should be 3 columns of data with the labels "subjID", "choice", "RT". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Choice made for the current trial, coded as 1/2 to indicate lower/upper boundary or left/right choices (e.g., 1 1 1 2 1 2).

RT Choice reaction time for the current trial, in ****seconds**** (e.g., 0.435 0.383 0.314 0.309, etc.).

**Note:* The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("choiceRT_ddm").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Ratcliff, R. (1978). A theory of memory retrieval. *Psychological Review*, 85(2), 59-108. <https://doi.org/10.1037/0033-295X.85.2.59>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- choiceRT_ddm(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- choiceRT_ddm(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

 choiceRT_ddm_single *Drift Diffusion Model*

Description

Individual Bayesian Modeling of the Choice Reaction Time Task using Drift Diffusion Model. It has the following parameters: alpha (boundary separation), beta (bias), delta (drift rate), tau (non-decision time).

- **Task:** Choice Reaction Time Task
- **Model:** Drift Diffusion Model (Ratcliff, 1978)

Usage

```
choiceRT_ddm_single(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "RT". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.

<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. Not available for this model.
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, it's possible to set model-specific argument(s) as follows: RTbound Floating point value representing the lower bound (i.e., minimum allowed) reaction time. Defaults to 0.1 (100 milliseconds).

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Choice Reaction Time Task, there should be 3 columns of data with the labels "subjID", "choice", "RT". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Choice made for the current trial, coded as 1/2 to indicate lower/upper boundary or left/right choices (e.g., 1 1 1 2 1 2).

RT Choice reaction time for the current trial, in ****seconds**** (e.g., 0.435 0.383 0.314 0.309, etc.).

**Note:* The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in

samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("choiceRT_ddm_single").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Ratcliff, R. (1978). A theory of memory retrieval. *Psychological Review*, 85(2), 59-108. <https://doi.org/10.1037/0033-295X.85.2.59>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- choiceRT_ddm_single(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- choiceRT_ddm_single(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

cra_exp

Exponential Subjective Value Model

Description

Hierarchical Bayesian Modeling of the Choice Under Risk and Ambiguity Task using Exponential Subjective Value Model. It has the following parameters: alpha (risk attitude), beta (ambiguity attitude), gamma (inverse temperature).

- **Task:** Choice Under Risk and Ambiguity Task
- **Model:** Exponential Subjective Value Model (Hsu et al., 2005)

Usage

```
cra_exp(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
```

```

    vb = FALSE,
    inc_postpred = FALSE,
    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "prob", "ambig", "reward_var", "reward_fix", "choice". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "sv", "sv_fix", "sv_var", "p_var".
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Choice Under Risk and Ambiguity Task, there should be 6 columns of data with the labels "subjID", "prob", "ambig", "reward_var", "reward_fix", "choice". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

prob Objective probability of the variable lottery.

ambig Ambiguity level of the variable lottery (0 for risky lottery; greater than 0 for ambiguous lottery).

reward_var Amount of reward in variable lottery. Assumed to be greater than zero.

reward_fix Amount of reward in fixed lottery. Assumed to be greater than zero.

choice If the variable lottery was selected, choice == 1; otherwise choice == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Jaeyeong Yang <<jaeyeong.yang1125@gmail.com>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("cra_exp").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Hsu, M., Bhatt, M., Adolphs, R., Tranel, D., & Camerer, C. F. (2005). Neural systems responding to degrees of uncertainty in human decision-making. *Science*, 310(5754), 1680-1683. <https://doi.org/10.1126/science.1115327>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- cra_exp(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- cra_exp(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
```

```
print_fit(output)

## End(Not run)
```

cra_linear	<i>Linear Subjective Value Model</i>
------------	--------------------------------------

Description

Hierarchical Bayesian Modeling of the Choice Under Risk and Ambiguity Task using Linear Subjective Value Model. It has the following parameters: alpha (risk attitude), beta (ambiguity attitude), gamma (inverse temperature).

- **Task:** Choice Under Risk and Ambiguity Task
- **Model:** Linear Subjective Value Model (Levy et al., 2010)

Usage

```
cra_linear(  
  data = NULL,  
  niter = 4000,  
  nwarmup = 1000,  
  nchain = 4,  
  ncore = 1,  
  nthin = 1,  
  inits = "vb",  
  ind_pars = "mean",  
  model_regressor = FALSE,  
  vb = FALSE,  
  inc_postpred = FALSE,  
  adapt_delta = 0.95,  
  stepsize = 1,  
  max_treedepth = 10,  
  seed = 42,  
  ...  
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "prob", "ambig", "reward_var", "reward_fix", "choice". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.

nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == nthin$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "sv", "sv_fix", "sv_var", "p_var".
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Choice Under Risk and Ambiguity Task, there should be 6 columns of data with the labels "subjID", "prob", "ambig", "reward_var", "reward_fix", "choice". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

prob Objective probability of the variable lottery.

ambig Ambiguity level of the variable lottery (0 for risky lottery; greater than 0 for ambiguous lottery).

reward_var Amount of reward in variable lottery. Assumed to be greater than zero.

reward_fix Amount of reward in fixed lottery. Assumed to be greater than zero.

choice If the variable lottery was selected, `choice == 1`; otherwise `choice == 0`.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Jaeyeong Yang <<jaeyeong.yang1125@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("cra_linear").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Levy, I., Snell, J., Nelson, A. J., Rustichini, A., & Glimcher, P. W. (2010). Neural representation of subjective value under risk and ambiguity. *Journal of Neurophysiology*, 103(2), 1036-1047.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- cra_linear(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- cra_linear(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

dbdm_prob_weight

Probability Weight Function

Description

Hierarchical Bayesian Modeling of the Description Based Decision Making Task using Probability Weight Function. It has the following parameters: tau (probability weight function), rho (subject utility function), lambda (loss aversion parameter), beta (inverse softmax temperature).

- **Task:** Description Based Decision Making Task
- **Model:** Probability Weight Function (Erev et al., 2010; Hertwig et al., 2004; Jessup et al., 2008)

Usage

```

dbdm_prob_weight(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjid", "opt1hprob", "opt2hprob", "opt1hval", "opt1lval", "opt2hval", "opt2lval", "choice". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every <code>i == nthin</code> sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Description Based Decision Making Task, there should be 8 columns of data with the labels "subjID", "opt1hprob", "opt2hprob", "opt1hval", "opt1lval", "opt2hval", "opt2lval", "choice". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

opt1hprob Possibility of getting higher value of outcome(opt1hval) when choosing option 1.

opt2hprob Possibility of getting higher value of outcome(opt2hval) when choosing option 2.

opt1hval Possible (with opt1hprob probability) outcome of option 1.

opt1lval Possible (with (1 - opt1hprob) probability) outcome of option 1.

opt2hval Possible (with opt2hprob probability) outcome of option 2.

opt2lval Possible (with (1 - opt2hprob) probability) outcome of option 2.

choice If option 1 was selected, choice == 1; else if option 2 was selected, choice == 2.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple

chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Yoonseo Zoh](#) <<zohyos7@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("`dbdm_prob_weight`").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

- Erev, I., Ert, E., Roth, A. E., Haruvy, E., Herzog, S. M., Hau, R., ... & Lebiere, C. (2010). A choice prediction competition: Choices from experience and from description. *Journal of Behavioral Decision Making*, 23(1), 15-47.
- Hertwig, R., Barron, G., Weber, E. U., & Erev, I. (2004). Decisions from experience and the effect of rare events in risky choice. *Psychological science*, 15(8), 534-539.
- Jessup, R. K., Bishara, A. J., & Busemeyer, J. R. (2008). Feedback produces divergence from prospect theory in descriptive choice. *Psychological Science*, 19(10), 1015-1022.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- dbdm_prob_weight(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- dbdm_prob_weight(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

dd_cs

Constant-Sensitivity (CS) Model

Description

Hierarchical Bayesian Modeling of the Delay Discounting Task using Constant-Sensitivity (CS) Model. It has the following parameters: r (exponential discounting rate), s (impatience), β (inverse temperature).

- **Task:** Delay Discounting Task
- **Model:** Constant-Sensitivity (CS) Model (Ebert et al., 2007)

Usage

```
dd_cs(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
```

```

    vb = FALSE,
    inc_postpred = FALSE,
    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Delay Discounting Task, there should be 6 columns of data with the labels "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

delay_later An integer representing the delayed days for the later option (e.g. 1, 6, 28).

amount_later A floating point number representing the amount for the later option (e.g. 10.5, 13.4, 30.9).

delay_sooner An integer representing the delayed days for the sooner option (e.g. 0).

amount_sooner A floating point number representing the amount for the sooner option (e.g. 10).

choice If amount_later was selected, choice == 1; else if amount_sooner was selected, choice == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan](#)

[User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("dd_cs").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Ebert, J. E. J., & Prelec, D. (2007). The Fragility of Time: Time-Insensitivity and Valuation of the Near and Far Future. *Management Science*. <https://doi.org/10.1287/mnsc.1060.0671>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- dd_cs(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- dd_cs(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)
```

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

dd_cs_single	<i>Constant-Sensitivity (CS) Model</i>
--------------	--

Description

Individual Bayesian Modeling of the Delay Discounting Task using Constant-Sensitivity (CS) Model. It has the following parameters: r (exponential discounting rate), s (impatience), β (inverse temperature).

- **Task:** Delay Discounting Task
- **Model:** Constant-Sensitivity (CS) Model (Ebert et al., 2007)

Usage

```
dd_cs_single(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.

<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Delay Discounting Task, there should be 6 columns of data with the labels "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

delay_later An integer representing the delayed days for the later option (e.g. 1, 6, 28).

amount_later A floating point number representing the amount for the later option (e.g. 10.5, 13.4, 30.9).

delay_sooner An integer representing the delayed days for the sooner option (e.g. 0).

amount_sooner A floating point number representing the amount for the sooner option (e.g. 10).

choice If amount_later was selected, choice == 1; else if amount_sooner was selected, choice == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The nwarmup argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("dd_cs_single").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Ebert, J. E. J., & Prelec, D. (2007). The Fragility of Time: Time-Insensitivity and Valuation of the Near and Far Future. *Management Science*. <https://doi.org/10.1287/mnsc.1060.0671>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- dd_cs_single(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- dd_cs_single(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

dd_exp

Exponential Model

Description

Hierarchical Bayesian Modeling of the Delay Discounting Task using Exponential Model. It has the following parameters: r (exponential discounting rate), β (inverse temperature).

- **Task:** Delay Discounting Task
- **Model:** Exponential Model (Samuelson, 1937)

Usage

```
dd_exp(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Delay Discounting Task, there should be 6 columns of data with the labels "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

delay_later An integer representing the delayed days for the later option (e.g. 1, 6, 28).

amount_later A floating point number representing the amount for the later option (e.g. 10.5, 13.4, 30.9).

delay_sooner An integer representing the delayed days for the sooner option (e.g. 0).

amount_sooner A floating point number representing the amount for the sooner option (e.g. 10).

choice If amount_later was selected, choice == 1; else if amount_sooner was selected, choice == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple

chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("dd_exp").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Samuelson, P. A. (1937). A Note on Measurement of Utility. The Review of Economic Studies, 4(2), 155. <https://doi.org/10.2307/2967612>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- dd_exp(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- dd_exp(
```

```

data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

dd_hyperbolic

Hyperbolic Model

Description

Hierarchical Bayesian Modeling of the Delay Discounting Task using Hyperbolic Model. It has the following parameters: k (discounting rate), beta (inverse temperature).

- **Task:** Delay Discounting Task
- **Model:** Hyperbolic Model (Mazur, 1987)

Usage

```

dd_hyperbolic(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Delay Discounting Task, there should be 6 columns of data with the labels "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

delay_later An integer representing the delayed days for the later option (e.g. 1, 6, 28).

amount_later A floating point number representing the amount for the later option (e.g. 10.5, 13.4, 30.9).

delay_sooner An integer representing the delayed days for the sooner option (e.g. 0).

amount_sooner A floating point number representing the amount for the sooner option (e.g. 10).

choice If amount_later was selected, choice == 1; else if amount_sooner was selected, choice == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The nwarmup argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("dd_hyperbolic").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Mazur, J. E. (1987). An adjustment procedure for studying delayed reinforcement.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- dd_hyperbolic(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- dd_hyperbolic(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

Description

Individual Bayesian Modeling of the Delay Discounting Task using Hyperbolic Model. It has the following parameters: k (discounting rate), β (inverse temperature).

- **Task:** Delay Discounting Task
- **Model:** Hyperbolic Model (Mazur, 1987)

Usage

```
dd_hyperbolic_single(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".

model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_tredepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Delay Discounting Task, there should be 6 columns of data with the labels "subjID", "delay_later", "amount_later", "delay_sooner", "amount_sooner", "choice". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

delay_later An integer representing the delayed days for the later option (e.g. 1, 6, 28).

amount_later A floating point number representing the amount for the later option (e.g. 10.5, 13.4, 30.9).

delay_sooner An integer representing the delayed days for the sooner option (e.g. 0).

amount_sooner A floating point number representing the amount for the sooner option (e.g. 10).

choice If amount_later was selected, choice == 1; else if amount_sooner was selected, choice == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains

begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("dd_hyperbolic_single").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Mazur, J. E. (1987). An adjustment procedure for studying delayed reinforcement.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- dd_hyperbolic_single(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- dd_hyperbolic_single(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

estimate_mode

Function to estimate mode of MCMC samples

Description

Based on codes from '<http://stackoverflow.com/questions/2547402/is-there-a-built-in-function-for-finding-the-mode>' see the comment by Rasmus Baath

Usage

```
estimate_mode(x)
```

Arguments

x MCMC samples or some numeric or array values.

extract_ic	<i>Extract Model Comparison Estimates</i>
------------	---

Description

Extract Model Comparison Estimates

Usage

```
extract_ic(model_data = NULL, ic = "looic", ncore = 2)
```

Arguments

model_data	Object returned by 'hBayesDM' model function
ic	Information Criterion. 'looic', 'waic', or 'both'
ncore	Number of cores to use when computing LOOIC

Value

IC Leave-One-Out and/or Watanabe-Akaike information criterion estimates.

Examples

```
## Not run:
library(hBayesDM)
output = bandit2arm_delta("example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 1)
extract_ic(output)
extract_ic(output, ic = "waic")

## End(Not run)
```

gng_ml	<i>RW + noise</i>
--------	-------------------

Description

Hierarchical Bayesian Modeling of the Orthogonalized Go/Nogo Task using RW + noise. It has the following parameters: xi (noise), ep (learning rate), rho (effective size).

- **Task:** Orthogonalized Go/Nogo Task
- **Model:** RW + noise (Guitart-Masip et al., 2012)

Usage

```

gng_m1(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "cue", "keyPressed", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "Qgo", "Qnogo", "Wgo", "Wnogo".
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_tredepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Orthogonalized Go/Nogo Task, there should be 4 columns of data with the labels "subjID", "cue", "keyPressed", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

cue Nominal integer representing the cue shown for that trial: 1, 2, 3, or 4.

keyPressed Binary value representing the subject's response for that trial (where Press == 1; No press == 0).

outcome Ternary value representing the outcome of that trial (where Positive feedback == 1; Neutral feedback == 0; Negative feedback == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("gng_m1").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Guitart-Masip, M., Huys, Q. J. M., Fuentemilla, L., Dayan, P., Duzel, E., & Dolan, R. J. (2012). Go and no-go learning in reward and punishment: Interactions between affect and effect. *Neuroimage*, 62(1), 154-166. <https://doi.org/10.1016/j.neuroimage.2012.04.024>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- gng_m1(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- gng_m1(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")
```

```

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

gng_m2	<i>RW + noise + bias</i>
--------	--------------------------

Description

Hierarchical Bayesian Modeling of the Orthogonalized Go/Nogo Task using RW + noise + bias. It has the following parameters: xi (noise), ep (learning rate), b (action bias), rho (effective size).

- **Task:** Orthogonalized Go/Nogo Task
- **Model:** RW + noise + bias (Guitart-Masip et al., 2012)

Usage

```

gng_m2(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "cue", "keyPressed", "outcome". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "Qgo", "Qnogo", "Wgo", "Wnogo".
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_tredepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Orthogonalized Go/Nogo Task, there should be 4 columns of data with the labels "subjID", "cue", "keyPressed", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

cue Nominal integer representing the cue shown for that trial: 1, 2, 3, or 4.

keyPressed Binary value representing the subject's response for that trial (where Press == 1; No press == 0).

outcome Ternary value representing the outcome of that trial (where Positive feedback == 1; Neutral feedback == 0; Negative feedback == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: **hBayesDM** 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("gng_m2").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Guitart-Masip, M., Huys, Q. J. M., Fuentemilla, L., Dayan, P., Duzel, E., & Dolan, R. J. (2012). Go and no-go learning in reward and punishment: Interactions between affect and effect. *Neuroimage*, 62(1), 154-166. <https://doi.org/10.1016/j.neuroimage.2012.04.024>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- gng_m2(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- gng_m2(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

Description

Hierarchical Bayesian Modeling of the Orthogonalized Go/Nogo Task using RW + noise + bias + pi. It has the following parameters: xi (noise), ep (learning rate), b (action bias), pi (Pavlovian bias), rho (effective size).

- **Task:** Orthogonalized Go/Nogo Task
- **Model:** RW + noise + bias + pi (Guitart-Masip et al., 2012)

Usage

```
gng_m3(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "cue", "keyPressed", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".

model_regressor	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "Qgo", "Qnogo", "Wgo", "Wnogo", "SV".
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Orthogonalized Go/Nogo Task, there should be 4 columns of data with the labels "subjID", "cue", "keyPressed", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

cue Nominal integer representing the cue shown for that trial: 1, 2, 3, or 4.

keyPressed Binary value representing the subject's response for that trial (where Press == 1; No press == 0).

outcome Ternary value representing the outcome of that trial (where Positive feedback == 1; Neutral feedback == 0; Negative feedback == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("gng_m3").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Guitart-Masip, M., Huys, Q. J. M., Fuentemilla, L., Dayan, P., Duzel, E., & Dolan, R. J. (2012). Go and no-go learning in reward and punishment: Interactions between affect and effect. *Neuroimage*, 62(1), 154-166. <https://doi.org/10.1016/j.neuroimage.2012.04.024>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- gng_m3(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- gng_m3(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

gng_m4

 $RW(\text{rew/pun}) + \text{noise} + \text{bias} + \pi$

Description

Hierarchical Bayesian Modeling of the Orthogonalized Go/Nogo Task using $RW(\text{rew/pun}) + \text{noise} + \text{bias} + \pi$. It has the following parameters: ξ (noise), ϵ (learning rate), b (action bias), π (Pavlovian bias), ρ_{rew} (reward sensitivity), ρ_{pun} (punishment sensitivity).

- **Task:** Orthogonalized Go/Nogo Task
- **Model:** $RW(\text{rew/pun}) + \text{noise} + \text{bias} + \pi$ (Cavanagh et al., 2013)

Usage

```
gng_m4(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
```

```

    vb = FALSE,
    inc_postpred = FALSE,
    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
  )

```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "cue", "keyPressed", "outcome". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every <code>i == nthin</code> sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "Qgo", "Qnogo", "Wgo", "Wnogo", "SV".
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Orthogonalized Go/Nogo Task, there should be 4 columns of data with the labels "subjID", "cue", "keyPressed", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

cue Nominal integer representing the cue shown for that trial: 1, 2, 3, or 4.

keyPressed Binary value representing the subject's response for that trial (where Press == 1; No press == 0).

outcome Ternary value representing the outcome of that trial (where Positive feedback == 1; Neutral feedback == 0; Negative feedback == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("gng_m4").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Cavanagh, J. F., Eisenberg, I., Guitart-Masip, M., Huys, Q., & Frank, M. J. (2013). Frontal Theta Overrides Pavlovian Learning Biases. *Journal of Neuroscience*, 33(19), 8541-8548. <https://doi.org/10.1523/JNEUROSCI.5751.2013>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- gng_m4(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- gng_m4(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

hdi	<i>Compute Highest-Density Interval</i>
-----	---

Description

Computes the highest density interval from a sample of representative values, estimated as shortest credible interval. Based on John Kruschke's codes.

Usage

```
hdi(sample_vec, ci_prob = 0.95)
```

Arguments

sample_vec	A vector of representative values from a probability distribution (e.g., MCMC samples).
ci_prob	A scalar between 0 and 1, indicating the mass within the credible interval that is to be estimated.

Value

A vector containing the limits of the HDI

hgf_ibrb	<i>Hierarchical Bayesian version of the Hierarchical Gaussian Filter model for binary inputs and binary responses</i>
----------	---

Description

Hierarchical Bayesian Modeling using Hierarchical Bayesian version of the Hierarchical Gaussian Filter model for binary inputs and binary responses. It has the following parameters: kappa (phasic volatility for coupling with higher level for each level ($2 \sim L-1$)), omega (tonic volatility for each level ($2 \sim L$)), zeta (inverse decision noise, the tendency to choose the response that corresponds with one's current belief).

- **Model:** Hierarchical Bayesian version of the Hierarchical Gaussian Filter model for binary inputs and binary responses (Mathys C, 2011; Mathys CD et al., 2014)

Usage

```

hgf_ibrb(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "trialNum", "u", "y". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. Not available for this model.

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, it's possible to set model-specific argument(s) as follows: L Total level of hierarchy. Defaults to minimum level of 3 input_first TRUE if participant observed <code>u[t]</code> before choosing <code>y[t]</code> , FALSE if participant observed <code>u[t]</code> after choosing <code>y[t]</code> mu0 prior belief for each level before starting the experiment sigma0 prior uncertainty for each level before starting the experiment kappa_lower Lower bounds for kappa for each level (2 ~ L-1). Defaults to [0] and can not be negative. Parameter value is fixed for level <code>l</code> if <code>kappa_upper[l] == kappa_lower[l]</code> . kappa_upper Upper bounds for kappa for each level (2 ~ L-1). Defaults to [2]. Parameter value is fixed for level <code>l</code> if <code>kappa_upper[l] == kappa_lower[l]</code> . omega_lower Lower bounds for omega for each level (2 ~ L). Defaults to [-10, -15]. Parameter value is fixed for level <code>l</code> if <code>omega_upper[l] == omega_lower[l]</code> . omega_upper Upper bounds for omega for each level (2 ~ L). Defaults to [0, 0]. Parameter value is fixed for level <code>l</code> if <code>omega_upper[l] == omega_lower[l]</code> . zeta_lower Upper bound for zeta. Defaults to 0 and can not be negative. Parameter value is fixed if <code>zeta_lower == zeta_upper</code> .

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Hierarchical Bayesian version of the Hierarchical Gaussian Filter model for binary inputs and binary responses, there should be 4 columns of data with the labels "subjID", "trialNum", "u", "y". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

trialNum Nominal integer representing the trial number: 1, 2, ...

u Integer value representing the input on that trial: 0 or 1.

y Integer value representing the subject's choice on that trial: 0 or 1.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the

necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Jinwoo Jeong <jwjeong96@gmail.com>](mailto:jwjeong96@gmail.com), [Juha Lee <juhajulia44@gmail.com>](mailto:juhajulia44@gmail.com), [Yusom Jo <yaun2288@snu.ac.kr>](mailto:yaun2288@snu.ac.kr)

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("hgf_ibrb").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

- Mathys C, Daunizeau J, Friston KJ and Stephan KE (2011) A Bayesian foundation for individual learning under uncertainty. *Front. Hum. Neurosci.* 5:39. <https://doi.org/10.3389/fnhum.2011.00039>
- Mathys CD, Lomakina EI, Daunizeau J, Iglesias S, Brodersen KH, Friston KJ and Stephan KE (2014) Uncertainty in perception and the Hierarchical Gaussian Filter. *Front. Hum. Neurosci.* 8:825. <https://doi.org/10.3389/fnhum.2014.00825>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- hgf_ibrb(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- hgf_ibrb(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

hgf_ibrb_single

Individual-level Bayesian version of the Hierarchical Gaussian Filter model for binary inputs and binary responses

Description

Individual Bayesian Modeling using Individual-level Bayesian version of the Hierarchical Gaussian Filter model for binary inputs and binary responses. It has the following parameters: kappa (phasic volatility for coupling with higher level for each level ($2 \sim L-1$)), omega (tonic volatility for each level ($2 \sim L$)), zeta (inverse decision noise, the tendency to choose the response that corresponds with one's current belief).

- **Model:** Individual-level Bayesian version of the Hierarchical Gaussian Filter model for binary inputs and binary responses (Mathys C, 2011; Mathys CD et al., 2014)

Usage

```
hgf_ibrb_single(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "trialNum", "u", "y". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.

<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. Not available for this model.
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, it's possible to set model-specific argument(s) as follows: L Total level of hierarchy. Defaults to minimum level of 3 input_first TRUE if participant observed <code>u[t]</code> before choosing <code>y[t]</code> , FALSE if participant observed <code>u[t]</code> after choosing <code>y[t]</code> mu0 prior belief for each level before starting the experiment sigma0 prior uncertainty for each level before starting the experiment kappa_lower Lower bounds for kappa for each level (2 ~ L-1). Defaults to [0] and can not be negative. Parameter value is fixed for level <code>l</code> if <code>kappa_upper[l] == kappa_lower[l]</code> . kappa_upper Upper bounds for kappa for each level (2 ~ L-1). Defaults to [2]. Parameter value is fixed for level <code>l</code> if <code>kappa_upper[l] == kappa_lower[l]</code> . omega_lower Lower bounds for omega for each level (2 ~ L). Defaults to [-10, -15]. Parameter value is fixed for level <code>l</code> if <code>omega_upper[l] == omega_lower[l]</code> . omega_upper Upper bounds for omega for each level (2 ~ L). Defaults to [0, 0]. Parameter value is fixed for level <code>l</code> if <code>omega_upper[l] == omega_lower[l]</code> . zeta_lower Upper bound for zeta. Defaults to 0 and can not be negative. Parameter value is fixed if <code>zeta_lower == zeta_upper</code> .

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Individual-level Bayesian version of the Hierarchical Gaussian Filter model for binary inputs and binary responses, there should be 3 columns of data with the labels "trialNum", "u", "y". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

trialNum Nominal integer representing the trial number: 1, 2, ...

u Integer value representing the input on that trial: 0 or 1.

y Integer value representing the subject's choice on that trial: 0 or 1.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Jinwoo Jeong <<jwjeong96@gmail.com>>, Juha Lee <<juhajulia44@gmail.com>>, Yusom Jo <<yaun2288@snu.ac.kr>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("hgf_ibrb_single").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

- Mathys C, Daunizeau J, Friston KJ and Stephan KE (2011) A Bayesian foundation for individual learning under uncertainty. *Front. Hum. Neurosci.* 5:39. <https://doi.org/10.3389/fnhum.2011.00039>
- Mathys CD, Lomakina EI, Daunizeau J, Iglesias S, Brodersen KH, Friston KJ and Stephan KE (2014) Uncertainty in perception and the Hierarchical Gaussian Filter. *Front. Hum. Neurosci.* 8:825. <https://doi.org/10.3389/fnhum.2014.00825>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- hgf_ibrb_single(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- hgf_ibrb_single(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

Description

Hierarchical Bayesian Modeling of the Iowa Gambling Task using Outcome-Representation Learning Model. It has the following parameters: Arew (reward learning rate), Apun (punishment learning rate), K (perseverance decay), betaF (outcome frequency weight), betaP (perseverance weight).

- **Task:** Iowa Gambling Task (Ahn et al., 2008)
- **Model:** Outcome-Representation Learning Model (Haines et al., 2018)

Usage

```

igt_orl(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, it's possible to set model-specific argument(s) as follows: payscale Raw payoffs within data are divided by this number. Used for scaling data. Defaults to 100.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Iowa Gambling Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer indicating which deck was chosen on that trial (where A==1, B==2, C==3, and D==4).

gain Floating point value representing the amount of currency won on that trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on that trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Nate Haines](#) <<haines.175@osu.edu>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("igt_orl").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

- Ahn, W. Y., Bussemeyer, J. R., & Wagenmakers, E. J. (2008). Comparison of decision learning models using the generalization criterion method. *Cognitive Science*, 32(8), 1376-1402. <https://doi.org/10.1080/03640210802352>
- Haines, N., Vassileva, J., & Ahn, W.-Y. (2018). The Outcome-Representation Learning Model: A Novel Reinforcement Learning Model of the Iowa Gambling Task. *Cognitive Science*. <https://doi.org/10.1111/cogs.12688>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- igt_orl(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)
```

```

# Run the model with example data
output <- igt_orl(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

igt_pvl_decay

Prospect Valence Learning (PVL) Decay-RI

Description

Hierarchical Bayesian Modeling of the Iowa Gambling Task using Prospect Valence Learning (PVL) Decay-RI. It has the following parameters: A (decay rate), alpha (outcome sensitivity), cons (response consistency), lambda (loss aversion).

- **Task:** Iowa Gambling Task (Ahn et al., 2008)
- **Model:** Prospect Valence Learning (PVL) Decay-RI (Ahn et al., 2014)

Usage

```

igt_pvl_decay(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every $i == nthin$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, it's possible to set model-specific argument(s) as follows: payscale Raw payoffs within data are divided by this number. Used for scaling data. Defaults to 100.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Iowa Gambling Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer indicating which deck was chosen on that trial (where A==1, B==2, C==3, and D==4).

gain Floating point value representing the amount of currency won on that trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on that trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("igt_pvl_decay").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Ahn, W. Y., Busemeyer, J. R., & Wagenmakers, E. J. (2008). Comparison of decision learning models using the generalization criterion method. *Cognitive Science*, 32(8), 1376-1402. <https://doi.org/10.1080/03640210802352>

Ahn, W.-Y., Vasilev, G., Lee, S.-H., Busemeyer, J. R., Kruschke, J. K., Bechara, A., & Vassileva, J. (2014). Decision-making in stimulant and opiate addicts in protracted abstinence: evidence from computational modeling with pure users. *Frontiers in Psychology*, 5, 1376. <https://doi.org/10.3389/fpsyg.2014.00849>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- igt_pvl_decay(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- igt_pvl_decay(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

Description

Hierarchical Bayesian Modeling of the Iowa Gambling Task using Prospect Valence Learning (PVL) Delta. It has the following parameters: A (learning rate), alpha (outcome sensitivity), cons (response consistency), lambda (loss aversion).

- **Task:** Iowa Gambling Task (Ahn et al., 2008)
- **Model:** Prospect Valence Learning (PVL) Delta (Ahn et al., 2008)

Usage

```
igt_pvl_delta(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".

model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: payscale Raw payoffs within data are divided by this number. Used for scaling data. Defaults to 100.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Iowa Gambling Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer indicating which deck was chosen on that trial (where A==1, B==2, C==3, and D==4).

gain Floating point value representing the amount of currency won on that trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on that trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("`igt_pvl_delta`").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Ahn, W. Y., Busemeyer, J. R., & Wagenmakers, E. J. (2008). Comparison of decision learning models using the generalization criterion method. *Cognitive Science*, 32(8), 1376-1402. <https://doi.org/10.1080/03640210802352>

Ahn, W. Y., Busemeyer, J. R., & Wagenmakers, E. J. (2008). Comparison of decision learning models using the generalization criterion method. *Cognitive Science*, 32(8), 1376-1402. <https://doi.org/10.1080/03640210802352>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- igt_pvl_delta(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- igt_pvl_delta(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

 igt_vpp

Value-Plus-Perseverance

Description

Hierarchical Bayesian Modeling of the Iowa Gambling Task using Value-Plus-Perseverance. It has the following parameters: A (learning rate), alpha (outcome sensitivity), cons (response consistency), lambda (loss aversion), epP (gain impact), epN (loss impact), K (decay rate), w (RL weight).

- **Task:** Iowa Gambling Task (Ahn et al., 2008)
- **Model:** Value-Plus-Perseverance (Worthy et al., 2013)

Usage

```
igt_vpp(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
```

```

    vb = FALSE,
    inc_postpred = FALSE,
    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
  )

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "gain", "loss". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: payscale Raw payoffs within data are divided by this number. Used for scaling data. Defaults to 100.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Iowa Gambling Task, there should be 4 columns of data with the labels "subjID", "choice", "gain", "loss". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer indicating which deck was chosen on that trial (where A==1, B==2, C==3, and D==4).

gain Floating point value representing the amount of currency won on that trial (e.g. 50, 100).

loss Floating point value representing the amount of currency lost on that trial (e.g. 0, -50).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("igt_vpp").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Ahn, W. Y., Busemeyer, J. R., & Wagenmakers, E. J. (2008). Comparison of decision learning models using the generalization criterion method. *Cognitive Science*, 32(8), 1376-1402. <https://doi.org/10.1080/03640210802352>

Worthy, D. A., & Todd Maddox, W. (2013). A comparison model of reinforcement-learning and win-stay-lose-shift decision-making processes: A tribute to W.K. Estes. *Journal of Mathematical Psychology*, 59, 41-49. <https://doi.org/10.1016/j.jmp.2013.10.001>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- igt_vpp(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- igt_vpp(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

multiplot	<i>Function to plot multiple figures</i>
-----------	--

Description

Plots multiple figures Based on codes from '[http://www.cookbook-r.com/Graphs/Multiple_graphs_on_one_page_\(ggplot2\)/](http://www.cookbook-r.com/Graphs/Multiple_graphs_on_one_page_(ggplot2)/)'

Usage

```
multiplot(..., plots = NULL, cols = NULL)
```

Arguments

...	Plot objects
plots	List containing plot objects
cols	Number of columns within the multi-figure plot

peer_ocu	<i>Other-Conferred Utility (OCU) Model</i>
----------	--

Description

Hierarchical Bayesian Modeling of the Peer Influence Task using Other-Conferred Utility (OCU) Model. It has the following parameters: rho (risk preference), tau (inverse temperature), ocu (other-conferred utility).

- **Task:** Peer Influence Task (Chung et al., 2015)
- **Model:** Other-Conferred Utility (OCU) Model

Usage

```
peer_ocu(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
```

```

    max_treedepth = 10,
    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "condition", "p_gamble", "safe_Hpayoff", "safe_Lpayoff", "risky_Hpayoff", "risky_Lpayoff", "choice". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for

the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Peer Influence Task, there should be 8 columns of data with the labels "subjID", "condition", "p_gamble", "safe_Hpayoff", "safe_Lpayoff", "risky_Hpayoff", "risky_Lpayoff", "choice". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

condition 0: solo, 1: info (safe/safe), 2: info (mix), 3: info (risky/risky).

p_gamble Probability of receiving a high payoff (same for both options).

safe_Hpayoff High payoff of the safe option.

safe_Lpayoff Low payoff of the safe option.

risky_Hpayoff High payoff of the risky option.

risky_Lpayoff Low payoff of the risky option.

choice Which option was chosen? 0: safe, 1: risky.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == nthin$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Harhim Park <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("peer_ocu").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Chung, D., Christopoulos, G. I., King-Casas, B., Ball, S. B., & Chiu, P. H. (2015). Social signals of safety and risk confer utility and have asymmetric effects on observers' choices. *Nature Neuroscience*, 18(6), 912-916.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- peer_ocu(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- peer_ocu(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

plot_dist	<i>Plots the histogram of MCMC samples.</i>
-----------	---

Description

Plots the histogram of MCMC samples.

Usage

```
plot_dist(  
  sample = NULL,  
  title = NULL,  
  x_lab = "Value",  
  y_lab = "Density",  
  x_lim = NULL,  
  font_size = NULL,  
  bin_size = NULL,  
  ...  
)
```

Arguments

sample	MCMC samples
title	Character value containing the main title for the plot
x_lab	Character value containing the x label
y_lab	Character value containing the y label
x_lim	Vector containing the lower and upper x-bounds of the plot
font_size	Size of the font to use for plotting. Defaults to 10
bin_size	Size of the bins for creating the histogram. Defaults to 30
...	Arguments that can be additionally supplied to geom_histogram

Value

h1 Plot object

plot_hdi	<i>Plots highest density interval (HDI) from (MCMC) samples and prints HDI in the R console. HDI is indicated by a red line. Based on John Kruschke's codes.</i>
----------	--

Description

Plots highest density interval (HDI) from (MCMC) samples and prints HDI in the R console. HDI is indicated by a red line. Based on John Kruschke's codes.

Usage

```
plot_hdi(  
  sample = NULL,  
  ci_prob = 0.95,  
  title = NULL,  
  x_lab = "Value",  
  y_lab = "Density",  
  font_size = NULL,  
  bin_size = 30,  
  ...  
)
```

Arguments

sample	MCMC samples
ci_prob	A scalar between 0 and 1, indicating the mass within the credible interval that is to be estimated.
title	Character value containing the main title for the plot
x_lab	Character value containing the x label
y_lab	Character value containing the y label
font_size	Integer value specifying the font size to be used for the plot labels
bin_size	Integer value specifying how wide the bars on the histogram should be. Defaults to 30.
...	Arguments that can be additionally supplied to geom_histogram

Value

A vector containing the limits of the HDI

plot_ind	<i>Plots individual posterior distributions using bayesplot.</i>
----------	---

Description

Plots individual posterior distributions using **bayesplot**.

Usage

```
plot_ind(obj = NULL, pars, show_density = T, ...)
```

Arguments

obj	An output of hBayesDM. Its class should be 'hBayesDM'.
pars	Character vector of parameter names to plot.
show_density	If TRUE, draws posterior densities via <code>bayesplot::mcmc_areas</code> ; otherwise draws point intervals via <code>bayesplot::mcmc_intervals</code> .
...	Additional arguments forwarded to the underlying bayesplot function.

Examples

```
## Not run:
# Run a model
output <- dd_hyperbolic("example", 2000, 1000, 3, 3)

# Plot the hyper parameters ('k' and 'beta')
plot(output)

# Plot individual 'k' (discounting rate) parameters
plot_ind(output, "k")

# Plot individual 'beta' (inverse temperature) parameters
plot_ind(output, "beta")

# Plot individual 'beta' parameters but don't show density
plot_ind(output, "beta", show_density = F)

## End(Not run)
```

print_fit	<i>Print model-fits (mean LOOIC or WAIC values in addition to Akaike weights) of hBayesDM Models</i>
-----------	--

Description

Print model-fits (mean LOOIC or WAIC values in addition to Akaike weights) of hBayesDM Models

Usage

```
print_fit(..., ic = "looic", ncore = 2, round_to = 3)
```

Arguments

...	Model objects output by hBayesDM functions (e.g. output1, output2, etc.)
ic	Which model comparison information criterion to use? 'looic', 'waic', or 'both'
ncore	Number of cores to use when computing LOOIC
round_to	Number of digits to the right of the decimal point in the output

Value

model_table A table with relevant model comparison data. LOOIC and WAIC weights are computed as Akaike weights.

Examples

```
## Not run:
# Run two models and store results in "output1" and "output2"
output1 <- dd_hyperbolic("example", 2000, 1000, 3, 3)

output2 <- dd_exp("example", 2000, 1000, 3, 3)

# Show the LOOIC model fit estimates
print_fit(output1, output2)

# To show the WAIC model fit estimates
print_fit(output1, output2, ic = "waic")

# To show both LOOIC and WAIC
print_fit(output1, output2, ic = "both")

## End(Not run)
```

prl_ewa

*Experience-Weighted Attraction Model***Description**

Hierarchical Bayesian Modeling of the Probabilistic Reversal Learning Task using Experience-Weighted Attraction Model. It has the following parameters: phi (1 - learning rate), rho (experience decay factor), beta (inverse temperature).

- **Task:** Probabilistic Reversal Learning Task
- **Model:** Experience-Weighted Attraction Model (Ouden et al., 2013)

Usage

```
prl_ewa(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.

<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "ev_c", "ev_nc", "ew_c", "ew_nc".
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Reversal Learning Task, there should be 3 columns of data with the labels "subjID", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on that trial: 1 or 2.

outcome Integer value representing the outcome of that trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Jaeyeong Yang \(for model-based regressors\)](#) <<jaeyeong.yang1125@gmail.com>>, [Harhim Park \(for model-based regressors\)](#) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("prl_ewa").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Ouden, den, H. E. M., Daw, N. D., Fernandez, G., Elshout, J. A., Rijpkema, M., Hoogman, M., et al. (2013). Dissociable Effects of Dopamine and Serotonin on Reversal Learning. *Neuron*, 80(4), 1090-1100. <https://doi.org/10.1016/j.neuron.2013.08.030>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- prl_ewa(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- prl_ewa(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

prl_fictitious

Fictitious Update Model

Description

Hierarchical Bayesian Modeling of the Probabilistic Reversal Learning Task using Fictitious Update Model. It has the following parameters: eta (learning rate), alpha (indecision point), beta (inverse temperature).

- **Task:** Probabilistic Reversal Learning Task
- **Model:** Fictitious Update Model (Glascher et al., 2009)

Usage

```
prl_fictitious(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
```

```

    vb = FALSE,
    inc_postpred = FALSE,
    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
  )

```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "outcome". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "ev_c", "ev_nc", "pe_c", "pe_nc", "dv".
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Reversal Learning Task, there should be 3 columns of data with the labels "subjID", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on that trial: 1 or 2.

outcome Integer value representing the outcome of that trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Jaeyeong Yang (for model-based regressors) <<jaeyeong.yang1125@gmail.com>>, Harhim Park (for model-based regressors) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("prl_fictitious").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Glascher, J., Hampton, A. N., & O'Doherty, J. P. (2009). Determining a Role for Ventromedial Prefrontal Cortex in Encoding Action-Based Value Signals During Reward-Related Decision Making. *Cerebral Cortex*, 19(2), 483-495. <https://doi.org/10.1093/cercor/bhn098>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- prl_fictitious(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- prl_fictitious(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

prl_fictitious_multipleB

Fictitious Update Model

Description

Multiple-Block Hierarchical Bayesian Modeling of the Probabilistic Reversal Learning Task using Fictitious Update Model. It has the following parameters: eta (learning rate), alpha (indecision point), beta (inverse temperature).

- **Task:** Probabilistic Reversal Learning Task
- **Model:** Fictitious Update Model (Glascher et al., 2009)

Usage

```
prl_fictitious_multipleB(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "block", "choice", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.

<code>nthin</code>	Every <code>i == nthin</code> sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "ev_c", "ev_nc", "pe_c", "pe_nc", "dv".
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Reversal Learning Task, there should be 4 columns of data with the labels "subjID", "block", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

block A unique identifier for each of the multiple blocks within each subject.

choice Integer value representing the option chosen on that trial: 1 or 2.

outcome Integer value representing the outcome of that trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Jaeyeong Yang \(for model-based regressors\)](#) <<jaeyeong.yang1125@gmail.com>>, [Harhim Park \(for model-based regressors\)](#) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("prl_fictitious_multipleB").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Glascher, J., Hampton, A. N., & O'Doherty, J. P. (2009). Determining a Role for Ventromedial Prefrontal Cortex in Encoding Action-Based Value Signals During Reward-Related Decision Making. *Cerebral Cortex*, 19(2), 483-495. <https://doi.org/10.1093/cercor/bhn098>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- prl_fictitious_multipleB(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- prl_fictitious_multipleB(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

prl_fictitious_rp	<i>Fictitious Update Model, with separate learning rates for positive and negative prediction error (PE)</i>
-------------------	--

Description

Hierarchical Bayesian Modeling of the Probabilistic Reversal Learning Task using Fictitious Update Model, with separate learning rates for positive and negative prediction error (PE). It has the following parameters: eta_pos (learning rate, +PE), eta_neg (learning rate, -PE), alpha (indecision point), beta (inverse temperature).

- **Task:** Probabilistic Reversal Learning Task
- **Model:** Fictitious Update Model, with separate learning rates for positive and negative prediction error (PE) (Glascher et al., 2009; Ouden et al., 2013)

Usage

```
prl_fictitious_rp(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "ev_c", "ev_nc", "pe_c", "pe_nc", "dv".
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Reversal Learning Task, there should be 3 columns of data with the labels "subjID", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on that trial: 1 or 2.

outcome Integer value representing the outcome of that trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users

change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Jaeyeong Yang (for model-based regressors) <<jaeyeong.yang1125@gmail.com>>, Harhim Park (for model-based regressors) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("prl_fictitious_rp").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Glascher, J., Hampton, A. N., & O'Doherty, J. P. (2009). Determining a Role for Ventromedial Prefrontal Cortex in Encoding Action-Based Value Signals During Reward-Related Decision Making. *Cerebral Cortex*, 19(2), 483-495. <https://doi.org/10.1093/cercor/bhn098>

Oudén, den, H. E. M., Daw, N. D., Fernandez, G., Elshout, J. A., Rijpkema, M., Hoogman, M., et al. (2013). Dissociable Effects of Dopamine and Serotonin on Reversal Learning. *Neuron*, 80(4), 1090-1100. <https://doi.org/10.1016/j.neuron.2013.08.030>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- prl_fictitious_rp(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
```

```

output <- prl_fictitious_rp(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

prl_fictitious_rp_woa *Fictitious Update Model, with separate learning rates for positive and negative prediction error (PE), without alpha (indecision point)*

Description

Hierarchical Bayesian Modeling of the Probabilistic Reversal Learning Task using Fictitious Update Model, with separate learning rates for positive and negative prediction error (PE), without alpha (indecision point). It has the following parameters: eta_pos (learning rate, +PE), eta_neg (learning rate, -PE), beta (inverse temperature).

- **Task:** Probabilistic Reversal Learning Task
- **Model:** Fictitious Update Model, with separate learning rates for positive and negative prediction error (PE), without alpha (indecision point) (Glascher et al., 2009; Ouden et al., 2013)

Usage

```

prl_fictitious_rp_woa(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,

```

```

    max_treedepth = 10,
    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "ev_c", "ev_nc", "pe_c", "pe_nc", "dv".
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for

the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Reversal Learning Task, there should be 3 columns of data with the labels "subjID", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on that trial: 1 or 2.

outcome Integer value representing the outcome of that trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Jaeyeong Yang (for model-based regressors) <<jaeyeong.yang1125@gmail.com>>, Harhim Park (for model-based regressors) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("prl_fictitious_rp_woa").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Glascher, J., Hampton, A. N., & O'Doherty, J. P. (2009). Determining a Role for Ventromedial Prefrontal Cortex in Encoding Action-Based Value Signals During Reward-Related Decision Making. *Cerebral Cortex*, 19(2), 483-495. <https://doi.org/10.1093/cercor/bhn098>

Ouden, den, H. E. M., Daw, N. D., Fernandez, G., Elshout, J. A., Rijpkema, M., Hoogman, M., et al. (2013). Dissociable Effects of Dopamine and Serotonin on Reversal Learning. *Neuron*, 80(4), 1090-1100. <https://doi.org/10.1016/j.neuron.2013.08.030>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- prl_fictitious_rp_woa(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- prl_fictitious_rp_woa(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

prl_fictitious_woa *Fictitious Update Model, without alpha (indecision point)*

Description

Hierarchical Bayesian Modeling of the Probabilistic Reversal Learning Task using Fictitious Update Model, without alpha (indecision point). It has the following parameters: eta (learning rate), beta (inverse temperature).

- **Task:** Probabilistic Reversal Learning Task
- **Model:** Fictitious Update Model, without alpha (indecision point) (Glascher et al., 2009)

Usage

```
prl_fictitious_woa(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.

<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "ev_c", "ev_nc", "pe_c", "pe_nc", "dv".
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Reversal Learning Task, there should be 3 columns of data with the labels "subjID", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on that trial: 1 or 2.

outcome Integer value representing the outcome of that trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Jaeyeong Yang \(for model-based regressors\)](#) <<jaeyeong.yang1125@gmail.com>>, [Harhim Park \(for model-based regressors\)](#) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("prl_fictitious_woa").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Glascher, J., Hampton, A. N., & O'Doherty, J. P. (2009). Determining a Role for Ventromedial Prefrontal Cortex in Encoding Action-Based Value Signals During Reward-Related Decision Making. *Cerebral Cortex*, 19(2), 483-495. <https://doi.org/10.1093/cercor/bhn098>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- prl_fictitious_woa(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- prl_fictitious_woa(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

prl_rp

Reward-Punishment Model

Description

Hierarchical Bayesian Modeling of the Probabilistic Reversal Learning Task using Reward-Punishment Model. It has the following parameters: Apun (punishment learning rate), Arew (reward learning rate), beta (inverse temperature).

- **Task:** Probabilistic Reversal Learning Task
- **Model:** Reward-Punishment Model (Ouden et al., 2013)

Usage

```
prl_rp(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
```

```

    vb = FALSE,
    inc_postpred = FALSE,
    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
  )

```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "outcome". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every <code>i == nthin</code> sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "ev_c", "ev_nc", "pe".
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Reversal Learning Task, there should be 3 columns of data with the labels "subjID", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value representing the option chosen on that trial: 1 or 2.

outcome Integer value representing the outcome of that trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Jaeyeong Yang (for model-based regressors) <<jaeyeong.yang1125@gmail.com>>, Harhim Park (for model-based regressors) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("prl_rp").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Ouden, den, H. E. M., Daw, N. D., Fernandez, G., Elshout, J. A., Rijpkema, M., Hoogman, M., et al. (2013). Dissociable Effects of Dopamine and Serotonin on Reversal Learning. *Neuron*, 80(4), 1090-1100. <https://doi.org/10.1016/j.neuron.2013.08.030>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- prl_rp(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- prl_rp(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

prl_rp_multipleB *Reward-Punishment Model*

Description

Multiple-Block Hierarchical Bayesian Modeling of the Probabilistic Reversal Learning Task using Reward-Punishment Model. It has the following parameters: Apun (punishment learning rate), Arew (reward learning rate), beta (inverse temperature).

- **Task:** Probabilistic Reversal Learning Task
- **Model:** Reward-Punishment Model (Ouden et al., 2013)

Usage

```
prl_rp_multipleB(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "block", "choice", "outcome". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.

<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "ev_c", "ev_nc", "pe".
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Reversal Learning Task, there should be 4 columns of data with the labels "subjID", "block", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

block A unique identifier for each of the multiple blocks within each subject.

choice Integer value representing the option chosen on that trial: 1 or 2.

outcome Integer value representing the outcome of that trial (where reward == 1, and loss == -1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument

can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Jaeyeong Yang (for model-based regressors) <<jaeyeong.yang1125@gmail.com>>, Harhim Park (for model-based regressors) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("prl_rp_multipleB").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Oudén, den, H. E. M., Daw, N. D., Fernandez, G., Elshout, J. A., Rijpkema, M., Hoogman, M., et al. (2013). Dissociable Effects of Dopamine and Serotonin on Reversal Learning. *Neuron*, 80(4), 1090-1100. <https://doi.org/10.1016/j.neuron.2013.08.030>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- prl_rp_multipleB(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- prl_rp_multipleB(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

pstRT_ddm

Drift Diffusion Model

Description

Hierarchical Bayesian Modeling of the Probabilistic Selection Task (with RT data) using Drift Diffusion Model. It has the following parameters: *a* (boundary separation), *tau* (non-decision time), *d1* (drift rate scaling), *d2* (drift rate scaling), *d3* (drift rate scaling).

- **Task:** Probabilistic Selection Task (with RT data) (Frank et al., 2007; Frank et al., 2004)
- **Model:** Drift Diffusion Model (Pedersen et al., 2017)

Usage

```
pstRT_ddm(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
```

```

ncore = 1,
nthin = 1,
inits = "vb",
ind_pars = "mean",
model_regressor = FALSE,
vb = FALSE,
inc_postpred = FALSE,
adapt_delta = 0.95,
stepsize = 1,
max_treedepth = 10,
seed = 42,
...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "cond", "choice", "RT". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == nthin$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "choice_os", "RT_os"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.

seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: RTbound Floating point value representing the lower bound (i.e., minimum allowed) reaction time. Defaults to 0.1 (100 milliseconds).

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Selection Task (with RT data), there should be 4 columns of data with the labels "subjID", "cond", "choice", "RT". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

cond Integer value representing the task condition of the given trial (AB == 1, CD == 2, EF == 3).

choice Integer value representing the option chosen on the given trial (1 or 2).

RT Float value representing the time taken for the response on the given trial.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler

control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Hoyoung Doh](mailto:hoyoung.doh@gmail.com) <<hoyoung.doh@gmail.com>>, [Sanghoon Kang](mailto:sanghoon.kang@yale.edu) <<sanghoon.kang@yale.edu>>, [Jihyun K. Hur](mailto:jihyun.hur@yale.edu) <<jihyun.hur@yale.edu>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("pstRT_ddm").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Frank, M. J., Santamaria, A., O'Reilly, R. C., & Willcutt, E. (2007). Testing computational models of dopamine and noradrenaline dysfunction in attention deficit/hyperactivity disorder. *Neuropsychopharmacology*, 32(7), 1583-1599.

Frank, M. J., Seeberger, L. C., & O'reilly, R. C. (2004). By carrot or by stick: cognitive reinforcement learning in parkinsonism. *Science*, 306(5703), 1940-1943.

Pedersen, M. L., Frank, M. J., & Biele, G. (2017). The drift diffusion model as the choice rule in reinforcement learning. *Psychonomic bulletin & review*, 24(4), 1234-1251.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- pstRT_ddm(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- pstRT_ddm(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
```

```

plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

pstRT_rlddm1

Reinforcement Learning Drift Diffusion Model 1

Description

Hierarchical Bayesian Modeling of the Probabilistic Selection Task (with RT data) using Reinforcement Learning Drift Diffusion Model 1. It has the following parameters: a (boundary separation), τ (non-decision time), v (drift rate scaling), α (learning rate).

- **Task:** Probabilistic Selection Task (with RT data) (Frank et al., 2007; Frank et al., 2004)
- **Model:** Reinforcement Learning Drift Diffusion Model 1 (Pedersen et al., 2017)

Usage

```

pstRT_rlddm1(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "cond", "prob", "choice", "RT", "feedback". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every <code>i == nthin</code> sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "Q1", "Q2".
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "choice_os", "RT_os", "choice_sm", "RT_sm", "fd_sm"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, it's possible to set model-specific argument(s) as follows: RTbound Floating point value representing the lower bound (i.e., minimum allowed) reaction time. Defaults to 0.1 (100 milliseconds). initQ Floating point value representing the model's initial value of any choice.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for

the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Selection Task (with RT data), there should be 6 columns of data with the labels "subjID", "cond", "prob", "choice", "RT", "feedback". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

cond Integer value representing the task condition of the given trial (AB == 1, CD == 2, EF == 3).

prob Float value representing the probability that a correct response (1) is rewarded in the current task condition.

choice Integer value representing the option chosen on the given trial (1 or 2).

RT Float value representing the time taken for the response on the given trial.

feedback Integer value representing the outcome of the given trial (where 'correct' == 1, and 'incorrect' == 0).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Hoyoung Doh <<hoyoung.doh@gmail.com>>, Sanghoon Kang <<sanghoon.kang@yale.edu>>, Jihyun K. Hur <<jihyun.hur@yale.edu>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("pstRT_rlddm1").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Frank, M. J., Santamaria, A., O'Reilly, R. C., & Willcutt, E. (2007). Testing computational models of dopamine and noradrenaline dysfunction in attention deficit/hyperactivity disorder. *Neuropsychopharmacology*, 32(7), 1583-1599.

Frank, M. J., Seeberger, L. C., & O'reilly, R. C. (2004). By carrot or by stick: cognitive reinforcement learning in parkinsonism. *Science*, 306(5703), 1940-1943.

Pedersen, M. L., Frank, M. J., & Biele, G. (2017). The drift diffusion model as the choice rule in reinforcement learning. *Psychonomic bulletin & review*, 24(4), 1234-1251.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- pstRT_rlddm1(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- pstRT_rlddm1(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
```

```

plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

pstRT_rlddm6

Reinforcement Learning Drift Diffusion Model 6

Description

Hierarchical Bayesian Modeling of the Probabilistic Selection Task (with RT data) using Reinforcement Learning Drift Diffusion Model 6. It has the following parameters: *a* (boundary separation), *bp* (boundary separation power), *tau* (non-decision time), *v* (drift rate scaling), *alpha_pos* (learning rate for positive prediction error), *alpha_neg* (learning rate for negative prediction error).

- **Task:** Probabilistic Selection Task (with RT data) (Frank et al., 2007; Frank et al., 2004)
- **Model:** Reinforcement Learning Drift Diffusion Model 6 (Pedersen et al., 2017)

Usage

```

pstRT_rlddm6(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

<i>data</i>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "iter", "cond", "prob", "choice", "RT", "feedback". See Details below for more information.
-------------	--

niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). For this model they are: "Q1", "Q2".
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "choice_os", "RT_os", "choice_sm", "RT_sm", "fd_sm"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: RTbound Floating point value representing the lower bound (i.e., minimum allowed) reaction time. Defaults to 0.1 (100 milliseconds). initQ Floating point value representing the model's initial value of any choice.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Selection Task (with RT data), there should be 7 columns of data with the labels "subjID", "iter", "cond", "prob", "choice", "RT", "feedback". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

iter Integer value representing the trial number for each task condition.

cond Integer value representing the task condition of the given trial (AB == 1, CD == 2, EF == 3).

prob Float value representing the probability that a correct response (1) is rewarded in the current task condition.

choice Integer value representing the option chosen on the given trial (1 or 2).

RT Float value representing the time taken for the response on the given trial.

feedback Integer value representing the outcome of the given trial (where 'correct' == 1, and 'incorrect' == 0).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Hoyoung Doh](#) <<hoyoung.doh@gmail.com>>, [Sanghoon Kang](#) <<sanghoon.kang@yale.edu>>, [Jihyun K. Hur](#) <<jihyun.hur@yale.edu>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("pstRT_rlddm6").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

model_regressor List object containing the extracted model-based regressors.

References

Frank, M. J., Santamaria, A., O'Reilly, R. C., & Willcutt, E. (2007). Testing computational models of dopamine and noradrenaline dysfunction in attention deficit/hyperactivity disorder. *Neuropsychopharmacology*, 32(7), 1583-1599.

Frank, M. J., Seeberger, L. C., & O'reilly, R. C. (2004). By carrot or by stick: cognitive reinforcement learning in parkinsonism. *Science*, 306(5703), 1940-1943.

Pedersen, M. L., Frank, M. J., & Biele, G. (2017). The drift diffusion model as the choice rule in reinforcement learning. *Psychonomic bulletin & review*, 24(4), 1234-1251.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- pstRT_rlddm6(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- pstRT_rlddm6(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
```

```
print_fit(output)

## End(Not run)
```

pst_gainloss_Q	<i>Gain-Loss Q Learning Model</i>
----------------	-----------------------------------

Description

Hierarchical Bayesian Modeling of the Probabilistic Selection Task using Gain-Loss Q Learning Model. It has the following parameters: alpha_pos (learning rate for positive feedbacks), alpha_neg (learning rate for negative feedbacks), beta (inverse temperature).

- **Task:** Probabilistic Selection Task
- **Model:** Gain-Loss Q Learning Model (Frank et al., 2007)

Usage

```
pst_gainloss_Q(  
  data = NULL,  
  niter = 4000,  
  nwarmup = 1000,  
  nchain = 4,  
  ncore = 1,  
  nthin = 1,  
  inits = "vb",  
  ind_pars = "mean",  
  model_regressor = FALSE,  
  vb = FALSE,  
  inc_postpred = FALSE,  
  adapt_delta = 0.95,  
  stepsize = 1,  
  max_treedepth = 10,  
  seed = 42,  
  ...  
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "type", "choice", "reward". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.

ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Selection Task, there should be 4 columns of data with the labels "subjID", "type", "choice", "reward". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

type Two-digit number indicating which pair of stimuli were presented for that trial, e.g. 12, 34, or 56. The digit on the left (tens-digit) indicates the presented stimulus for option1, while the digit on the right (ones-digit) indicates that for option2. Code for each stimulus type (1~6) is defined as for 80% (type 1), 20% (type 2), 70% (type 3), 30% (type 4), 60% (type 5), 40% (type 6). The modeling will still work even if different probabilities are used for the stimuli; however, the total number of stimuli should be less than or equal to 6.

choice Whether the subject chose the left option (option1) out of the given two options (i.e. if option1 was chosen, 1; if option2 was chosen, 0).

reward Amount of reward earned as a result of the trial.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Jaeyeong Yang](#) <<jaeyeong.yang1125@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("pst_gainloss_Q").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Frank, M. J., Moustafa, A. A., Haughey, H. M., Curran, T., & Hutchison, K. E. (2007). Genetic triple dissociation reveals multiple roles for dopamine in reinforcement learning. *Proceedings of the National Academy of Sciences*, 104(41), 16311-16316.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- pst_gainloss_Q(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- pst_gainloss_Q(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

pst_Q

Q Learning Model

Description

Hierarchical Bayesian Modeling of the Probabilistic Selection Task using Q Learning Model. It has the following parameters: alpha (learning rate), beta (inverse temperature).

- **Task:** Probabilistic Selection Task
- **Model:** Q Learning Model (Frank et al., 2007)

Usage

```

pst_Q(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "type", "choice", "reward". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"

<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Probabilistic Selection Task, there should be 4 columns of data with the labels "subjID", "type", "choice", "reward". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

type Two-digit number indicating which pair of stimuli were presented for that trial, e.g. 12, 34, or 56. The digit on the left (tens-digit) indicates the presented stimulus for option1, while the digit on the right (ones-digit) indicates that for option2. Code for each stimulus type (1~6) is defined as for 80% (type 1), 20% (type 2), 70% (type 3), 30% (type 4), 60% (type 5), 40% (type 6). The modeling will still work even if different probabilities are used for the stimuli; however, the total number of stimuli should be less than or equal to 6.

choice Whether the subject chose the left option (option1) out of the given two options (i.e. if option1 was chosen, 1; if option2 was chosen, 0).

reward Amount of reward earned as a result of the trial.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple

chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: `hBayesDM` 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [David Munoz Tord](#) <<david.munoztord@unige.ch>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("pst_Q").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A `CmdStanMCMC` object (or `CmdStanVB` when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Frank, M. J., Moustafa, A. A., Haughey, H. M., Curran, T., & Hutchison, K. E. (2007). Genetic triple dissociation reveals multiple roles for dopamine in reinforcement learning. *Proceedings of the National Academy of Sciences*, 104(41), 16311-16316.

See Also

We refer users to our in-depth tutorial for an example of using `hBayesDM`: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- pst_Q(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)
```

```

# Run the model with example data
output <- pst_Q(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

ra_noLA

Prospect Theory, without loss aversion (LA) parameter

Description

Hierarchical Bayesian Modeling of the Risk Aversion Task using Prospect Theory, without loss aversion (LA) parameter. It has the following parameters: rho (risk aversion), tau (inverse temperature).

- **Task:** Risk Aversion Task
- **Model:** Prospect Theory, without loss aversion (LA) parameter (Sokol-Hessner et al., 2009)

Usage

```

ra_noLA(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,

```

```

    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "gain", "loss", "cert", "gamble". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial

observations and columns represent variables.

For the Risk Aversion Task, there should be 5 columns of data with the labels "subjID", "gain", "loss", "cert", "gamble". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

gain Possible (50%) gain outcome of a risky option (e.g. 9).

loss Possible (50%) loss outcome of a risky option (e.g. 5, or -5).

cert Guaranteed amount of a safe option. "cert" is assumed to be zero or greater than zero.

gamble If gamble was taken, gamble == 1; else gamble == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("ra_noLA").

- all_ind_pars** Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.
- par_vals** List object containing the posterior samples over different parameters.
- fit** A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.
- raw_data** Data.frame containing the raw data used to fit the model, as specified by the user.

References

Sokol-Hessner, P., Hsu, M., Curley, N. G., Delgado, M. R., Camerer, C. F., Phelps, E. A., & Smith, E. E. (2009). Thinking like a Trader Selectively Reduces Individuals' Loss Aversion. *Proceedings of the National Academy of Sciences of the United States of America*, 106(13), 5035-5040. <https://www.pnas.org/content/106/13/5035>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- ra_noLA(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- ra_noLA(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

ra_noRA	<i>Prospect Theory, without risk aversion (RA) parameter</i>
---------	--

Description

Hierarchical Bayesian Modeling of the Risk Aversion Task using Prospect Theory, without risk aversion (RA) parameter. It has the following parameters: lambda (loss aversion), tau (inverse temperature).

- **Task:** Risk Aversion Task
- **Model:** Prospect Theory, without risk aversion (RA) parameter (Sokol-Hessner et al., 2009)

Usage

```
ra_noRA(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "gain", "loss", "cert", "gamble". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every i == nthin sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.

<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Risk Aversion Task, there should be 5 columns of data with the labels "subjID", "gain", "loss", "cert", "gamble". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

gain Possible (50%) gain outcome of a risky option (e.g. 9).

loss Possible (50%) loss outcome of a risky option (e.g. 5, or -5).

cert Guaranteed amount of a safe option. "cert" is assumed to be zero or greater than zero.

gamble If gamble was taken, gamble == 1; else gamble == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains

begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: `hBayesDM` 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("ra_noRA").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A `CmdStanMCMC` object (or `CmdStanVB` when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Sokol-Hessner, P., Hsu, M., Curley, N. G., Delgado, M. R., Camerer, C. F., Phelps, E. A., & Smith, E. E. (2009). Thinking like a Trader Selectively Reduces Individuals' Loss Aversion. *Proceedings of the National Academy of Sciences of the United States of America*, 106(13), 5035-5040. <https://www.pnas.org/content/106/13/5035>

See Also

We refer users to our in-depth tutorial for an example of using `hBayesDM`: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- ra_noRA(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- ra_noRA(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

ra_prospect

Prospect Theory

Description

Hierarchical Bayesian Modeling of the Risk Aversion Task using Prospect Theory. It has the following parameters: rho (risk aversion), lambda (loss aversion), tau (inverse temperature).

- **Task:** Risk Aversion Task
- **Model:** Prospect Theory (Sokol-Hessner et al., 2009)

Usage

```
ra_prospect(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
```

```

    inc_postpred = FALSE,
    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "gain", "loss", "cert", "gamble". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == nthin$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Risk Aversion Task, there should be 5 columns of data with the labels "subjID", "gain", "loss", "cert", "gamble". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

gain Possible (50%) gain outcome of a risky option (e.g. 9).

loss Possible (50%) loss outcome of a risky option (e.g. 5, or -5).

cert Guaranteed amount of a safe option. "cert" is assumed to be zero or greater than zero.

gamble If gamble was taken, gamble == 1; else gamble == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("ra_prospect").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Sokol-Hessner, P., Hsu, M., Curley, N. G., Delgado, M. R., Camerer, C. F., Phelps, E. A., & Smith, E. E. (2009). Thinking like a Trader Selectively Reduces Individuals' Loss Aversion. *Proceedings of the National Academy of Sciences of the United States of America*, 106(13), 5035-5040. <https://www.pnas.org/content/106/13/5035>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- ra_prospect(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- ra_prospect(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

rdt_happiness

*Happiness Computational Model***Description**

Hierarchical Bayesian Modeling of the Risky Decision Task using Happiness Computational Model. It has the following parameters: w_0 (baseline), w_1 (weight of certain rewards), w_2 (weight of expected values), w_3 (weight of reward prediction errors), gam (forgetting factor), sig (standard deviation of error).

- **Task:** Risky Decision Task
- **Model:** Happiness Computational Model (Rutledge et al., 2014)

Usage

```
rdt_happiness(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "gain", "loss", "cert", "type", "gamble", "outcome", "happy", "RT_happy". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.

<code>nthin</code>	Every <code>i == nthin</code> sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Risky Decision Task, there should be 9 columns of data with the labels "subjID", "gain", "loss", "cert", "type", "gamble", "outcome", "happy", "RT_happy". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

gain Possible (50%) gain outcome of a risky option (e.g. 9).

loss Possible (50%) loss outcome of a risky option (e.g. 5, or -5).

cert Guaranteed amount of a safe option.

type `loss == -1`, `mixed == 0`, `gain == 1`

gamble If gamble was taken, `gamble == 1`; else `gamble == 0`.

outcome Result of the trial.

happy Happiness score.

RT_happy Reaction time for answering the happiness score.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Harhim Park](#) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("rdt_happiness").

all_ind_pars Data.frame containing the summarized parameter values (as specified by **ind_pars**) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when **vb = TRUE**) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Rutledge, R. B., Skandali, N., Dayan, P., & Dolan, R. J. (2014). A computational and neural model of momentary subjective well-being. *Proceedings of the National Academy of Sciences*, 111(33), 12252-12257.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- rdt_happiness(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- rdt_happiness(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

rhat

Function for extracting Rhat values from an hBayesDM object

Description

A convenience function for extracting Rhat values from an hBayesDM object. Can also check if all Rhat values are less than or equal to a specified value. If variational inference was used, an error message will be displayed.

Usage

```
rhat(fit = NULL, less = NULL)
```

Arguments

<code>fit</code>	Model output of class hBayesDM
<code>less</code>	A numeric value specifying how to check Rhat values. Defaults to FALSE.

Value

If 'less' is specified, then `rhat(fit, less)` will return TRUE if all Rhat values are less than or equal to 'less'. If any values are greater than 'less', `rhat(fit, less)` will return FALSE. If 'less' is left unspecified (NULL), `rhat(fit)` will return a data.frame object containing all Rhat values.

<code>task2AFC_sdt</code>	<i>Signal detection theory model</i>
---------------------------	--------------------------------------

Description

Hierarchical Bayesian Modeling of the 2-alternative forced choice task using Signal detection theory model. It has the following parameters: d (discriminability), c (decision bias (criteria)).

- **Task:** 2-alternative forced choice task
- **Model:** Signal detection theory model

Usage

```
task2AFC_sdt(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "stimulus", "response". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_tredepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the 2-alternative forced choice task, there should be 3 columns of data with the labels "subjID", "stimulus", "response". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

stimulus Types of Stimuli (Should be 1 or 0. 1 for Signal and 0 for Noise)

response Types of Responses (It should be same format as the stimulus field. Should be 1 or 0)

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Heesun Park](#) <<heesunpark26@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("task2AFC_sdt").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- task2AFC_sdt(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- task2AFC_sdt(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

ts_par4

Hybrid Model, with 4 parameters

Description

Hierarchical Bayesian Modeling of the Two-Step Task using Hybrid Model, with 4 parameters. It has the following parameters: α (learning rate for both stages 1 & 2), β (inverse temperature for both stages 1 & 2), π (perseverance), w (model-based weight).

- **Task:** Two-Step Task (Daw et al., 2011)
- **Model:** Hybrid Model, with 4 parameters (Daw et al., 2011; Wunderlich et al., 2012)

Usage

```
ts_par4(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
```

```

ncore = 1,
nthin = 1,
inits = "vb",
ind_pars = "mean",
model_regressor = FALSE,
vb = FALSE,
inc_postpred = FALSE,
adapt_delta = 0.95,
stepsize = 1,
max_treedepth = 10,
seed = 42,
...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "sub-jID", "level1_choice", "level2_choice", "reward". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == nthin$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred_step1", "y_pred_step2"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.

max_treepdepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: trans_prob Common state transition probability from Stage (Level) 1 to Stage (Level) 2. Defaults to 0.7.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Two-Step Task, there should be 4 columns of data with the labels "subjID", "level1_choice", "level2_choice", "reward". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

level1_choice Choice made for Level (Stage) 1 (1: stimulus 1, 2: stimulus 2).

level2_choice Choice made for Level (Stage) 2 (1: stimulus 3, 2: stimulus 4, 3: stimulus 5, 4: stimulus 6).

Note that, in our notation, choosing stimulus 1 in Level 1 leads to stimulus 3 & 4 in Level 2 with a common (0.7 by default) transition. Similarly, choosing stimulus 2 in Level 1 leads to stimulus 5 & 6 in Level 2 with a common (0.7 by default) transition. To change this default transition probability, set the function argument 'trans_prob' to your preferred value.

reward Reward after Level 2 (0 or 1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The nwarmup argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every $i == \text{nthin}$ samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: **adapt_delta**, **stepsize**, and **max_treedepth** are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Harhim Park](#) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("ts_par4").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Daw, N. D., Gershman, S. J., Seymour, B., Ben Seymour, Dayan, P., & Dolan, R. J. (2011). Model-Based Influences on Humans' Choices and Striatal Prediction Errors. *Neuron*, 69(6), 1204-1215. <https://doi.org/10.1016/j.neuron.2011.02.027>

Daw, N. D., Gershman, S. J., Seymour, B., Ben Seymour, Dayan, P., & Dolan, R. J. (2011). Model-Based Influences on Humans' Choices and Striatal Prediction Errors. *Neuron*, 69(6), 1204-1215. <https://doi.org/10.1016/j.neuron.2011.02.027>

Wunderlich, K., Smittenaar, P., & Dolan, R. J. (2012). Dopamine enhances model-based over model-free choice behavior. *Neuron*, 75(3), 418-424.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- ts_par4(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- ts_par4(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

ts_par6

Hybrid Model, with 6 parameters

Description

Hierarchical Bayesian Modeling of the Two-Step Task using Hybrid Model, with 6 parameters. It has the following parameters: a1 (learning rate in stage 1), beta1 (inverse temperature in stage 1), a2 (learning rate in stage 2), beta2 (inverse temperature in stage 2), pi (perseverance), w (model-based weight).

- **Task:** Two-Step Task (Daw et al., 2011)
- **Model:** Hybrid Model, with 6 parameters (Daw et al., 2011)

Usage

```
ts_par6(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
```

```

    model_regressor = FALSE,
    vb = FALSE,
    inc_postpred = FALSE,
    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "level1_choice", "level2_choice", "reward". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred_step1", "y_pred_step2"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.

... For this model, it's possible to set **model-specific argument(s)** as follows:

trans_prob Common state transition probability from Stage (Level) 1 to Stage (Level) 2. Defaults to 0.7.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Two-Step Task, there should be 4 columns of data with the labels "subjID", "level1_choice", "level2_choice", "reward". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

level1_choice Choice made for Level (Stage) 1 (1: stimulus 1, 2: stimulus 2).

level2_choice Choice made for Level (Stage) 2 (1: stimulus 3, 2: stimulus 4, 3: stimulus 5, 4: stimulus 6).

Note that, in our notation, choosing stimulus 1 in Level 1 leads to stimulus 3 & 4 in Level 2 with a common (0.7 by default) transition. Similarly, choosing stimulus 2 in Level 1 leads to stimulus 5 & 6 in Level 2 with a common (0.7 by default) transition. To change this default transition probability, set the function argument 'trans_prob' to your preferred value.

reward Reward after Level 2 (0 or 1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to

'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Harhim Park](#) <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("ts_par6").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Daw, N. D., Gershman, S. J., Seymour, B., Ben Seymour, Dayan, P., & Dolan, R. J. (2011). Model-Based Influences on Humans' Choices and Striatal Prediction Errors. *Neuron*, 69(6), 1204-1215. <https://doi.org/10.1016/j.neuron.2011.02.027>

Daw, N. D., Gershman, S. J., Seymour, B., Ben Seymour, Dayan, P., & Dolan, R. J. (2011). Model-Based Influences on Humans' Choices and Striatal Prediction Errors. *Neuron*, 69(6), 1204-1215. <https://doi.org/10.1016/j.neuron.2011.02.027>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- ts_par6(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- ts_par6(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
```

```

plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)

```

ts_par7

Hybrid Model, with 7 parameters (original model)

Description

Hierarchical Bayesian Modeling of the Two-Step Task using Hybrid Model, with 7 parameters (original model). It has the following parameters: a1 (learning rate in stage 1), beta1 (inverse temperature in stage 1), a2 (learning rate in stage 2), beta2 (inverse temperature in stage 2), pi (perseverance), w (model-based weight), lambda (eligibility trace).

- **Task:** Two-Step Task (Daw et al., 2011)
- **Model:** Hybrid Model, with 7 parameters (original model) (Daw et al., 2011)

Usage

```

ts_par7(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-seperated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "level1_choice", "level2_choice", "reward". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred_step1", "y_pred_step2"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, it's possible to set model-specific argument(s) as follows: trans_prob Common state transition probability from Stage (Level) 1 to Stage (Level) 2. Defaults to 0.7.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial

observations and columns represent variables.

For the Two-Step Task, there should be 4 columns of data with the labels "subjID", "level1_choice", "level2_choice", "reward". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

level1_choice Choice made for Level (Stage) 1 (1: stimulus 1, 2: stimulus 2).

level2_choice Choice made for Level (Stage) 2 (1: stimulus 3, 2: stimulus 4, 3: stimulus 5, 4: stimulus 6).

Note that, in our notation, choosing stimulus 1 in Level 1 leads to stimulus 3 & 4 in Level 2 with a common (0.7 by default) transition. Similarly, choosing stimulus 2 in Level 1 leads to stimulus 5 & 6 in Level 2 with a common (0.7 by default) transition. To change this default transition probability, set the function argument 'trans_prob' to your preferred value.

reward Reward after Level 2 (0 or 1).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: Harhim Park <<hrpark12@gmail.com>>

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("ts_par7").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Daw, N. D., Gershman, S. J., Seymour, B., Ben Seymour, Dayan, P., & Dolan, R. J. (2011). Model-Based Influences on Humans' Choices and Striatal Prediction Errors. *Neuron*, 69(6), 1204-1215. <https://doi.org/10.1016/j.neuron.2011.02.027>

Daw, N. D., Gershman, S. J., Seymour, B., Ben Seymour, Dayan, P., & Dolan, R. J. (2011). Model-Based Influences on Humans' Choices and Striatal Prediction Errors. *Neuron*, 69(6), 1204-1215. <https://doi.org/10.1016/j.neuron.2011.02.027>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- ts_par7(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- ts_par7(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

ug_bayes

*Ideal Observer Model***Description**

Hierarchical Bayesian Modeling of the Norm-Training Ultimatum Game using Ideal Observer Model. It has the following parameters: alpha (envy), beta (guilt), tau (inverse temperature).

- **Task:** Norm-Training Ultimatum Game
- **Model:** Ideal Observer Model (Xiang et al., 2013)

Usage

```
ug_bayes(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "offer", "accept". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.

<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Norm-Training Ultimatum Game, there should be 3 columns of data with the labels "subjID", "offer", "accept". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

offer Floating point value representing the offer made in that trial (e.g. 4, 10, 11).

accept 1 or 0, indicating whether the offer was accepted in that trial (where accepted == 1, rejected == 0).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument

can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("ug_bayes").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Xiang, T., Lohrenz, T., & Montague, P. R. (2013). Computational Substrates of Norms and Their Violations during Social Exchange. *Journal of Neuroscience*, 33(3), 1099-1108. <https://doi.org/10.1523/JNEUROSCI.1642-12.2013>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- ug_bayes(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- ug_bayes(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

ug_delta

Rescorla-Wagner (Delta) Model

Description

Hierarchical Bayesian Modeling of the Norm-Training Ultimatum Game using Rescorla-Wagner (Delta) Model. It has the following parameters: alpha (envy), tau (inverse temperature), ep (norm adaptation rate).

- **Task:** Norm-Training Ultimatum Game
- **Model:** Rescorla-Wagner (Delta) Model (Gu et al., 2015)

Usage

```
ug_delta(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
```

```

    vb = FALSE,
    inc_postpred = FALSE,
    adapt_delta = 0.95,
    stepsize = 1,
    max_treedepth = 10,
    seed = 42,
    ...
)

```

Arguments

data	Data to be modeled. It should be given as a data.frame object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "offer", "accept". See Details below for more information.
niter	Number of iterations, including warm-up. Defaults to 4000.
nwarmup	Number of iterations used for warm-up only. Defaults to 1000.
nchain	Number of Markov chains to run. Defaults to 4.
ncore	Number of CPUs to be used for running. Defaults to 1.
nthin	Every $i == \text{nthin}$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.
inits	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
ind_pars	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
model_regressor	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
vb	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
inc_postpred	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
adapt_delta	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
stepsize	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
max_treedepth	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
seed	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
...	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Norm-Training Ultimatum Game, there should be 3 columns of data with the labels "subjID", "offer", "accept". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

offer Floating point value representing the offer made in that trial (e.g. 4, 10, 11).

accept 1 or 0, indicating whether the offer was accepted in that trial (where accepted == 1, rejected == 0).

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The **nwarmup** argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, **nthin** is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Value

A class "hBayesDM" object model_data with the following components:

model Character value that is the name of the model ("ug_delta").

all_ind_pars Data.frame containing the summarized parameter values (as specified by ind_pars) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when vb = TRUE) produced by **cmdstanr**. Use fit\$draws(), fit\$summary(), etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Gu, X., Wang, X., Hula, A., Wang, S., Xu, S., Lohrenz, T. M., et al. (2015). Necessary, Yet Dissociable Contributions of the Insular and Ventromedial Prefrontal Cortices to Norm Adaptation: Computational and Lesion Evidence in Humans. *Journal of Neuroscience*, 35(2), 467-473. <https://doi.org/10.1523/JNEUROSCI.2906-14.2015>

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- ug_delta(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- ug_delta(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

wcs_sql

*Sequential Learning Model***Description**

Hierarchical Bayesian Modeling of the Wisconsin Card Sorting Task using Sequential Learning Model. It has the following parameters: r (reward sensitivity), p (punishment sensitivity), d (decision consistency or inverse temperature).

- **Task:** Wisconsin Card Sorting Task
- **Model:** Sequential Learning Model (Bishara et al., 2010)

Usage

```
wcs_sql(
  data = NULL,
  niter = 4000,
  nwarmup = 1000,
  nchain = 4,
  ncore = 1,
  nthin = 1,
  inits = "vb",
  ind_pars = "mean",
  model_regressor = FALSE,
  vb = FALSE,
  inc_postpred = FALSE,
  adapt_delta = 0.95,
  stepsize = 1,
  max_treedepth = 10,
  seed = 42,
  ...
)
```

Arguments

<code>data</code>	Data to be modeled. It should be given as a <code>data.frame</code> object, a filepath for a tab-separated txt file, "example" to use example data, or "choose" to choose data with an interactive window. Columns in the dataset must include: "subjID", "choice", "outcome". See Details below for more information.
<code>niter</code>	Number of iterations, including warm-up. Defaults to 4000.
<code>nwarmup</code>	Number of iterations used for warm-up only. Defaults to 1000.
<code>nchain</code>	Number of Markov chains to run. Defaults to 4.
<code>ncore</code>	Number of CPUs to be used for running. Defaults to 1.
<code>nthin</code>	Every $i == nthin$ sample will be used to generate the posterior distribution. Defaults to 1. A higher number can be used when auto-correlation within the MCMC sampling is high.

<code>inits</code>	Character value specifying how the initial values should be generated. Possible options are "vb" (default), "fixed", "random", or your own initial values.
<code>ind_pars</code>	Character value specifying how to summarize individual parameters. Current options are: "mean", "median", or "mode".
<code>model_regressor</code>	Whether to export model-based regressors (TRUE or FALSE). Not available for this model.
<code>vb</code>	Use variational inference to approximately draw from a posterior distribution. Defaults to FALSE.
<code>inc_postpred</code>	Include trial-level posterior predictive simulations in model output (may greatly increase file size). Defaults to FALSE. If set to TRUE, it includes: "y_pred"
<code>adapt_delta</code>	Floating point value representing the target acceptance probability of a new sample in the MCMC chain. Must be between 0 and 1. See Details below.
<code>stepsize</code>	Integer value specifying the size of each leapfrog step that the MCMC sampler can take on each new iteration. See Details below.
<code>max_treedepth</code>	Integer value specifying how many leapfrog steps the MCMC sampler can take on each new iteration. See Details below.
<code>seed</code>	Integer seed used for MCMC sampling and (when applicable) variational inference, to make results reproducible. Defaults to 42.
<code>...</code>	For this model, there is no model-specific argument.

Details

This section describes some of the function arguments in greater detail.

data should be assigned a character value specifying the full path and name (including extension information, e.g. ".txt") of the file that contains the behavioral data-set of all subjects of interest for the current analysis. The file should be a **tab-delimited** text file, whose rows represent trial-by-trial observations and columns represent variables.

For the Wisconsin Card Sorting Task, there should be 3 columns of data with the labels "subjID", "choice", "outcome". It is not necessary for the columns to be in this particular order, however it is necessary that they be labeled correctly and contain the information below:

subjID A unique identifier for each subject in the data-set.

choice Integer value indicating which deck was chosen on that trial: 1, 2, 3, or 4.

outcome 1 or 0, indicating the outcome of that trial: correct == 1, wrong == 0.

*Note: The file may contain other columns of data (e.g. "ReactionTime", "trial_number", etc.), but only the data within the column names listed above will be used during the modeling. As long as the necessary columns mentioned above are present and labeled correctly, there is no need to remove other miscellaneous data columns.

nwarmup is a numerical value that specifies how many MCMC samples should not be stored upon the beginning of each chain. For those familiar with Bayesian methods, this is equivalent to burn-in samples. Due to the nature of the MCMC algorithm, initial values (i.e. where the sampling chains begin) can have a heavy influence on the generated posterior distributions. The `nwarmup` argument can be set to a high number in order to curb the effects that initial values have on the resulting posteriors.

nchain is a numerical value that specifies how many chains (i.e. independent sampling sequences) should be used to draw samples from the posterior distribution. Since the posteriors are generated from a sampling process, it is good practice to run multiple chains to ensure that a reasonably representative posterior is attained. When the sampling is complete, it is possible to check the multiple chains for convergence by running the following line of code: `plot(output, type = "trace")`. The trace-plot should resemble a "furry caterpillar".

nthin is a numerical value that specifies the "skipping" behavior of the MCMC sampler, using only every `i == nthin` samples to generate posterior distributions. By default, `nthin` is equal to 1, meaning that every sample is used to generate the posterior.

Control Parameters: `adapt_delta`, `stepsize`, and `max_treedepth` are advanced options that give the user more control over Stan's MCMC sampler. It is recommended that only advanced users change the default values, as alterations can profoundly change the sampler's behavior. Refer to 'The No-U-Turn Sampler: Adaptively Setting Path Lengths in Hamiltonian Monte Carlo (Hoffman & Gelman, 2014, Journal of Machine Learning Research)' for more information on the sampler control parameters. One can also refer to 'Section 34.2. HMC Algorithm Parameters' of the [Stan User's Guide and Reference Manual](#), or to the help pages for `cmdstanr::sample()` for a less technical description of these arguments.

First-fit compile cost: hBayesDM 2.0 compiles each Stan model on first use via **cmdstanr** (no install-time precompile). Expect a one-time ~30s wait per model; subsequent fits reuse the cached binary.

Contributors: [Dayeong Min](#) <<mindy2801@snu.ac.kr>>

Value

A class "hBayesDM" object `model_data` with the following components:

model Character value that is the name of the model ("wcs_sql").

all_ind_pars Data.frame containing the summarized parameter values (as specified by `ind_pars`) for each subject.

par_vals List object containing the posterior samples over different parameters.

fit A CmdStanMCMC object (or CmdStanVB when `vb = TRUE`) produced by **cmdstanr**. Use `fit$draws()`, `fit$summary()`, etc. to interact with the posterior; see **cmdstanr** documentation for details.

raw_data Data.frame containing the raw data used to fit the model, as specified by the user.

References

Bishara, A. J., Kruschke, J. K., Stout, J. C., Bechara, A., McCabe, D. P., & Busemeyer, J. R. (2010). Sequential learning models for the Wisconsin card sort task: Assessing processes in substance dependent individuals. *Journal of Mathematical Psychology*, 54(1), 5-13.

See Also

We refer users to our in-depth tutorial for an example of using hBayesDM: <https://rpubs.com/CCSL/hBayesDM>

Examples

```
## Not run:
# Run the model with a given data.frame as df
output <- wcs_sql(
  data = df, niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Run the model with example data
output <- wcs_sql(
  data = "example", niter = 2000, nwarmup = 1000, nchain = 4, ncore = 4)

# Visually check convergence of the sampling chains (should look like 'hairy caterpillars')
plot(output, type = "trace")

# Check Rhat values (all Rhat values should be less than or equal to 1.1)
rhat(output)

# Plot the posterior distributions of the hyper-parameters (distributions should be unimodal)
plot(output)

# Show the WAIC and LOOIC model fit estimates
print_fit(output)

## End(Not run)
```

Index

alt_delta, 4
alt_gamma, 7

bandit2arm_delta, 11
bandit4arm2_kalman_filter, 14
bandit4arm_2par_lapse, 18
bandit4arm_4par, 21
bandit4arm_lapse, 25
bandit4arm_lapse_decay, 28
bandit4arm_singleA_lapse, 32
banditNarm_2par_lapse, 35
banditNarm_4par, 39
banditNarm_delta, 42
banditNarm_kalman_filter, 46
banditNarm_lapse, 49
banditNarm_lapse_decay, 53
banditNarm_singleA_lapse, 57
bart_ewmv, 60
bart_par4, 64

cgt_cm, 67
choiceRT_ddm, 71
choiceRT_ddm_single, 75
cra_exp, 78
cra_linear, 82

dbdm_prob_weight, 85
dd_cs, 89
dd_cs_single, 93
dd_exp, 96
dd_hyperbolic, 100
dd_hyperbolic_single, 103

estimate_mode, 107
extract_ic, 108

gng_m1, 108
gng_m2, 112
gng_m3, 115
gng_m4, 119

hdi, 123
hgf_ibrb, 123
hgf_ibrb_single, 127

igt_orl, 131
igt_pvl_decay, 135
igt_pvl_delta, 138
igt_vpp, 142

multiplot, 146

peer_ocu, 146
plot_dist, 150
plot_hdi, 151
plot_ind, 152
print_fit, 153
prl_ewa, 154
prl_fictitious, 157
prl_fictitious_multipleB, 161
prl_fictitious_rp, 164
prl_fictitious_rp_woa, 168
prl_fictitious_woa, 172
prl_rp, 175
prl_rp_multipleB, 179
pst_gainloss_Q, 194
pst_Q, 197
pstRT_ddm, 182
pstRT_rlddm1, 186
pstRT_rlddm6, 190

ra_noLA, 201
ra_noRA, 205
ra_prospect, 208
rdt_happiness, 212
rhat, 215

task2AFC_sdt, 216
ts_par4, 219
ts_par6, 223
ts_par7, 227

ug_bayes, [231](#)

ug_delta, [234](#)

wcs_sql, [238](#)