

Package: loo (via r-universe)

July 24, 2026

Type Package

Title Efficient Leave-One-Out Cross-Validation and WAIC for Bayesian Models

Version 2.10.1

Date 2026-07-23

Maintainer Jonah Gabry <jgabry@gmail.com>

Description Efficient approximate leave-one-out cross-validation (LOO) for Bayesian models fit using Markov chain Monte Carlo, as described in Vehtari, Gelman, and Gabry (2017) <[doi:10.1007/s11222-016-9696-4](https://doi.org/10.1007/s11222-016-9696-4)>. The approximation uses Pareto smoothed importance sampling (PSIS), a new procedure for regularizing importance weights. As a byproduct of the calculations, we also obtain approximate standard errors for estimated predictive errors and for the comparison of predictive errors between models. The package also provides methods for using stacking and other model weighting techniques to average Bayesian predictive distributions.

License GPL (>=3)

URL <https://mc-stan.org/loo/>, <https://discourse.mc-stan.org>

BugReports <https://github.com/stan-dev/loo/issues>

Depends R (>= 3.5)

Imports checkmate, matrixStats (>= 0.52), parallel, posterior (>= 1.7.0), stats

Suggests bayesplot (>= 1.7.0), brms (>= 2.10.0), ggplot2, graphics, knitr, rmarkdown, rstan, rstanarm (>= 2.19.0), rstantools, spdep, testthat (>= 3.0)

VignetteBuilder knitr

Config/testthat/edition 3

Config/testthat/parallel true

Config/testthat/start-first loo_subsampling_cases, loo_subsampling

Encoding UTF-8
LazyData TRUE
Roxygen list(markdown = TRUE)
SystemRequirements pandoc (>= 1.12.3), pandoc-citeproc
Config/roxygen2/version 8.0.0
Repository <https://r-multiverse-staging.r-universe.dev>
Date/Publication 2026-07-24 01:47:08 UTC
RemoteUrl <https://github.com/stan-dev/loo>
RemoteRef v2.10.1
RemoteSha df27ade95093db66cebf19d55530439386ac1a97

Contents

loo-package	3
ap_psis	5
compare	6
crps	7
E_loo	10
elpd	12
example_loglik_array	13
extract_log_lik	14
find_model_names	15
gpdfit	15
importance_sampling	16
kfold-generic	18
kfold-helpers	18
loo	20
loo-datasets	25
loo-glossary	26
loo_approximate_posterior	30
loo_compare	33
loo_model_weights	35
loo_moment_match	40
loo_moment_match_split	42
loo_predictive_metric	44
loo_subsample	45
nobs.psis_loo_ss	49
obs_idx	49
pareto-k-diagnostic	50
pointwise	53
print.loo	54
psis	55
psislw	57
relative_eff	59
sis	61

tis	63
update.psis_loo_ss	65
waic	67
weights.importance_sampling	69

Index	71
--------------	-----------

loo-package	<i>Efficient LOO-CV and WAIC for Bayesian models</i>
-------------	--

Description

Stan Development Team

This package implements the methods described in Vehtari, Gelman, and Gabry (2017), Vehtari, Simpson, Gelman, Yao, and Gabry (2024), and Yao et al. (2018). To get started see the **loo** package [vignettes](#), the `loo()` function for efficient approximate leave-one-out cross-validation (LOO-CV), the `psis()` function for the Pareto smoothed importance sampling (PSIS) algorithm, or `loo_model_weights()` for an implementation of Bayesian stacking of predictive distributions from multiple models.

Details

Leave-one-out cross-validation (LOO-CV) and the widely applicable information criterion (WAIC) are methods for estimating pointwise out-of-sample prediction accuracy from a fitted Bayesian model using the log-likelihood evaluated at the posterior simulations of the parameter values. LOO-CV and WAIC have various advantages over simpler estimates of predictive error such as AIC and DIC but are less used in practice because they involve additional computational steps. This package implements the fast and stable computations for approximate LOO-CV laid out in Vehtari, Gelman, and Gabry (2017). From existing posterior simulation draws, we compute LOO-CV using Pareto smoothed importance sampling (PSIS; Vehtari, Simpson, Gelman, Yao, and Gabry, 2024), a new procedure for stabilizing and diagnosing importance weights. As a byproduct of our calculations, we also obtain approximate standard errors for estimated predictive errors and for comparing of predictive errors between two models.

We recommend PSIS-LOO-CV instead of WAIC, because PSIS provides useful diagnostics and effective sample size and Monte Carlo standard error estimates.

Author(s)

Maintainer: Jonah Gabry <jgabry@gmail.com>

Authors:

- Jonah Gabry <jgabry@gmail.com>
- Aki Vehtari <Aki.Vehtari@aalto.fi>
- Måns Magnusson
- Yuling Yao
- Paul-Christian Bürkner
- Topi Paananen

- Andrew Gelman

Other contributors:

- Ben Goodrich [contributor]
- Juho Piironen [contributor]
- Bruno Nicenboim [contributor]
- Leevi Lindgren [contributor]
- Visruth Srimath Kandali [contributor]
- Florence Bockting [contributor]

References

- Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).
- Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)
- Yao, Y., Vehtari, A., Simpson, D., and Gelman, A. (2018) Using stacking to average Bayesian predictive distributions. *Bayesian Analysis*, advance publication, doi:10.1214/17-BA1091. ([online](#)).
- Magnusson, M., Riis Andersen, M., Jonasson, J. and Vehtari, A. (2019). Leave-One-Out Cross-Validation for Large Data. In *Thirty-sixth International Conference on Machine Learning*, PMLR 97:4244-4253.
- Magnusson, M., Riis Andersen, M., Jonasson, J. and Vehtari, A. (2020). Leave-One-Out Cross-Validation for Model Comparison in Large Data. In *Proceedings of the 23rd International Conference on Artificial Intelligence and Statistics (AISTATS)*, PMLR 108:341-351.
- Epifani, I., MacEachern, S. N., and Peruggia, M. (2008). Case-deletion importance sampling estimators: Central limit theorems and related results. *Electronic Journal of Statistics* 2, 774-806.
- Gelfand, A. E. (1996). Model determination using sampling-based methods. In *Markov Chain Monte Carlo in Practice*, ed. W. R. Gilks, S. Richardson, D. J. Spiegelhalter, 145-162. London: Chapman and Hall.
- Gelfand, A. E., Dey, D. K., and Chang, H. (1992). Model determination using predictive distributions with implementation via sampling-based methods. In *Bayesian Statistics 4*, ed. J. M. Bernardo, J. O. Berger, A. P. Dawid, and A. F. M. Smith, 147-167. Oxford University Press.
- Gelman, A., Hwang, J., and Vehtari, A. (2014). Understanding predictive information criteria for Bayesian models. *Statistics and Computing* 24, 997-1016.
- Ionides, E. L. (2008). Truncated importance sampling. *Journal of Computational and Graphical Statistics* 17, 295-311.
- Koopman, S. J., Shephard, N., and Creal, D. (2009). Testing the assumptions behind importance sampling. *Journal of Econometrics* 149, 2-11.
- Peruggia, M. (1997). On the variability of case-deletion importance sampling weights in the Bayesian linear model. *Journal of the American Statistical Association* 92, 199-207.
- Stan Development Team (2017). The Stan C++ Library, Version 2.17.0. <https://mc-stan.org>.

Stan Development Team (2018). RStan: the R interface to Stan, Version 2.17.3. <https://mc-stan.org>.

Watanabe, S. (2010). Asymptotic equivalence of Bayes cross validation and widely applicable information criterion in singular learning theory. *Journal of Machine Learning Research* **11**, 3571-3594.

Zhang, J., and Stephens, M. A. (2009). A new and efficient estimation method for the generalized Pareto distribution. *Technometrics* **51**, 316-325.

See Also

Useful links:

- <https://mc-stan.org/loo/>
- <https://discourse.mc-stan.org>
- Report bugs at <https://github.com/stan-dev/loo/issues>

ap_psis	<i>Pareto smoothed importance sampling (PSIS) using approximate posteriors</i>
---------	--

Description

Pareto smoothed importance sampling (PSIS) using approximate posteriors

Usage

```
ap_psis(log_ratios, log_p, log_g, ...)

## S3 method for class 'array'
ap_psis(log_ratios, log_p, log_g, ..., cores = getOption("mc.cores", 1))

## S3 method for class 'matrix'
ap_psis(log_ratios, log_p, log_g, ..., cores = getOption("mc.cores", 1))

## Default S3 method:
ap_psis(log_ratios, log_p, log_g, ...)
```

Arguments

log_ratios	The log-likelihood ratios (ie -log_lik)
log_p	The log-posterior (target) evaluated at S samples from the proposal distribution (g). A vector of length S.
log_g	The log-density (proposal) evaluated at S samples from the proposal distribution (g). A vector of length S.
...	Currently not in use.

- `cores` The number of cores to use for parallelization. This defaults to the option `mc.cores` which can be set for an entire R session by `options(mc.cores = NUMBER)`. The old option `loo.cores` is now deprecated but will be given precedence over `mc.cores` until `loo.cores` is removed in a future release. **As of version 2.0.0 the default is now 1 core if `mc.cores` is not set**, but we recommend using as many (or close to as many) cores as possible.
- Note for Windows 10 users: it is **strongly recommended** to avoid using the `.Rprofile` file to set `mc.cores` (using the `cores` argument or setting `mc.cores` interactively or in a script is fine).

Methods (by class)

- `ap_psis(array)`: An I by C by N array, where I is the number of MCMC iterations per chain, C is the number of chains, and N is the number of data points.
- `ap_psis(matrix)`: An S by N matrix, where S is the size of the posterior sample (with all chains merged) and N is the number of data points.
- `ap_psis(default)`: A vector of length S (posterior sample size).

`compare`

Model comparison (deprecated, old version)

Description

This function is deprecated. Please use the new `loo_compare()` function instead.

Usage

```
compare(..., x = list())
```

Arguments

- `...` At least two objects returned by `loo()` (or `waic()`).
- `x` A list of at least two objects returned by `loo()` (or `waic()`). This argument can be used as an alternative to specifying the objects in `...`

Details

When comparing two fitted models, we can estimate the difference in their expected predictive accuracy by the difference in `elpd_loo` or `elpd_waic` (or multiplied by -2, if desired, to be on the deviance scale).

When that difference, `elpd_diff`, is positive then the expected predictive accuracy for the second model is higher. A negative `elpd_diff` favors the first model.

When using `compare()` with more than two models, the values in the `elpd_diff` and `se_diff` columns of the returned matrix are computed by making pairwise comparisons between each model and the model with the best ELPD (i.e., the model in the first row). Although the `elpd_diff`

column is equal to the difference in `elpd_loo`, do not expect the `se_diff` column to be equal to the difference in `se_elpd_loo`.

To compute the standard error of the difference in ELPD we use a paired estimate to take advantage of the fact that the same set of N data points was used to fit both models. These calculations should be most useful when N is large, because then non-normality of the distribution is not such an issue when estimating the uncertainty in these sums. These standard errors, for all their flaws, should give a better sense of uncertainty than what is obtained using the current standard approach of comparing differences of deviances to a Chi-squared distribution, a practice derived for Gaussian linear models or asymptotically, and which only applies to nested models in any case.

Value

A vector or matrix with class `'compare.loo'` that has its own print method. If exactly two objects are provided in `...` or `x`, then the difference in expected predictive accuracy and the standard error of the difference are returned. If more than two objects are provided then a matrix of summary information is returned (see **Details**).

References

Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).

Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)

Examples

```
## Not run:
loo1 <- loo(log_lik1)
loo2 <- loo(log_lik2)
print(compare(loo1, loo2), digits = 3)
print(compare(x = list(loo1, loo2)))

waic1 <- waic(log_lik1)
waic2 <- waic(log_lik2)
compare(waic1, waic2)

## End(Not run)
```

Description

The `crps()` and `scrps()` functions and their `loo_*`() counterparts can be used to compute the continuously ranked probability score (CRPS) and scaled CRPS (SCRPS) (as defined by Bolin and Wallin, 2023). CRPS is a proper scoring rule, and strictly proper when the first moment of the predictive distribution is finite. Both can be expressed in terms of samples from the predictive distribution. See, for example, a paper by Gneiting and Raftery (2007) for a comprehensive discussion on CRPS.

Usage

```
crps(x, ...)  
  
scrps(x, ...)  
  
loo_crps(x, ...)  
  
loo_scrps(x, ...)  
  
## S3 method for class 'matrix'  
crps(x, x2, y, ..., permutations = 1)  
  
## S3 method for class 'numeric'  
crps(x, x2, y, ..., permutations = 1)  
  
## S3 method for class 'matrix'  
loo_crps(  
  x,  
  x2,  
  y,  
  log_lik,  
  ...,  
  permutations = 1,  
  r_eff = 1,  
  cores = getOption("mc.cores", 1)  
)  
  
## S3 method for class 'matrix'  
scrps(x, x2, y, ..., permutations = 1)  
  
## S3 method for class 'numeric'  
scrps(x, x2, y, ..., permutations = 1)  
  
## S3 method for class 'matrix'  
loo_scrps(  
  x,  
  x2,  
  y,  
  log_lik,
```



```

...,
permutations = 1,
r_eff = 1,
cores = getOption("mc.cores", 1)
)

```

Arguments

<code>x</code>	A S by N matrix (draws by observations), or a vector of length S when only single observation is provided in <code>y</code> .
<code>...</code>	Passed on to <code>E_loo()</code> in the <code>loo_*()</code> version of these functions.
<code>x2</code>	Independent draws from the same distribution as draws in <code>x</code> . Should be of the identical dimension.
<code>y</code>	A vector of observations or a single value.
<code>permutations</code>	An integer, with default value of 1, specifying how many times the expected value of $ X - X' $ ($ x - x2 $) is computed. The row order of <code>x2</code> is shuffled as elements <code>x</code> and <code>x2</code> are typically drawn given the same values of parameters. This happens, e.g., when one calls <code>posterior_predict()</code> twice for a fitted rstanarm or brms model. Generating more permutations is expected to decrease the variance of the computed expected value.
<code>log_lik</code>	A log-likelihood matrix the same size as <code>x</code> .
<code>r_eff</code>	An optional vector of relative effective sample size estimates containing one element per observation. See <code>psis()</code> for details.
<code>cores</code>	The number of cores to use for parallelization of <code>[psis()]</code> . See <code>psis()</code> for details.

Details

To compute (S)CRPS, the user needs to provide two sets of draws, `x` and `x2`, from the predictive distribution. This is due to the fact that formulas used to compute CRPS involve an expectation of the absolute difference of `x` and `x2`, both having the same distribution. See the `permutations` argument, as well as Gneiting and Raftery (2007) for details.

Value

A list containing two elements: `estimates` and `pointwise`. The former reports estimator and standard error and latter the pointwise values. Following Bolin & Wallin (2023), a larger value is better.

References

- Bolin, D., & Wallin, J. (2023). Local scale invariance and robustness of proper scoring rules. *Statistical Science*, 38(1):140-159.
- Gneiting, T., & Raftery, A. E. (2007). Strictly Proper Scoring Rules, Prediction, and Estimation. *Journal of the American Statistical Association*, 102(477), 359–378.

Examples

```
## Not run:
# An example using rstanarm
library(rstanarm)
data("kidiq")
fit <- stan_glm(kid_score ~ mom_hs + mom_iq, data = kidiq)
ypred1 <- posterior_predict(fit)
ypred2 <- posterior_predict(fit)
crps(ypred1, ypred2, y = fit$y)
loo_crps(ypred1, ypred2, y = fit$y, log_lik = log_lik(fit))

## End(Not run)
```

E_loo

Compute weighted expectations

Description

The `E_loo()` function computes weighted expectations (means, variances, quantiles) using the importance weights obtained from the [PSIS](#) smoothing procedure. The expectations estimated by the `E_loo()` function assume that the PSIS approximation is working well. **A small [Pareto k](#) estimate is necessary, but not sufficient, for `E_loo()` to give reliable estimates.** If the `log_ratios` argument is provided, `E_loo()` also computes a function specific Pareto k diagnostic, which must also be small for a reliable estimate. See more details below.

Usage

```
E_loo(x, psis_object, ...)

## Default S3 method:
E_loo(
  x,
  psis_object,
  ...,
  type = c("mean", "variance", "sd", "quantile"),
  probs = NULL,
  log_ratios = NULL
)

## S3 method for class 'matrix'
E_loo(
  x,
  psis_object,
  ...,
  type = c("mean", "variance", "sd", "quantile"),
  probs = NULL,
  log_ratios = NULL
)
```

Arguments

<code>x</code>	A numeric vector or matrix.
<code>psis_object</code>	An object returned by <code>psis()</code> .
<code>...</code>	Arguments passed to individual methods.
<code>type</code>	The type of expectation to compute. The options are "mean", "variance", "sd", and "quantile".
<code>probs</code>	For computing quantiles, a vector of probabilities.
<code>log_ratios</code>	Optionally, a vector or matrix (the same dimensions as <code>x</code>) of raw (not smoothed) log ratios. If working with log-likelihood values, the log ratios are the negative of those values. If <code>log_ratios</code> is specified we are able to compute more accurate Pareto k diagnostics specific to <code>E_loo()</code> .

Value

A named list with the following components:

`value` The result of the computation.

For the matrix method, `value` is a vector with `ncol(x)` elements, with one exception: when `type="quantile"` and multiple values are specified in `probs` the `value` component of the returned object is a `length(probs)` by `ncol(x)` matrix.

For the default/vector method the `value` component is scalar, with one exception: when `type="quantile"` and multiple values are specified in `probs` the `value` component is a vector with `length(probs)` elements.

`pareto_k` Function-specific diagnostic.

For the matrix method it will be a vector of length `ncol(x)` containing estimates of the shape parameter k of the generalized Pareto distribution. For the default/vector method, the estimate is a scalar. If `log_ratios` is not specified when calling `E_loo()`, the smoothed log-weights are used to estimate Pareto- k 's, which may produce optimistic estimates.

For `type="mean"`, `type="var"`, and `type="sd"`, the returned Pareto- k is usually the maximum of the Pareto- k 's for the left and right tail of hr and the right tail of r , where r is the importance ratio and $h = x$ for `type="mean"` and $h = x^2$ for `type="var"` and `type="sd"`. If h is binary, constant, or not finite, or if `type="quantile"`, the returned Pareto- k is the Pareto- k for the right tail of r .

Examples

```
if (requireNamespace("rstanarm", quietly = TRUE)) {
  # Use rstanarm package to quickly fit a model and get both a log-likelihood
  # matrix and draws from the posterior predictive distribution
  library("rstanarm")

  # data from help("lm")
  ctl <- c(4.17, 5.58, 5.18, 6.11, 4.50, 4.61, 5.17, 4.53, 5.33, 5.14)
  trt <- c(4.81, 4.17, 4.41, 3.59, 5.87, 3.83, 6.03, 4.89, 4.32, 4.69)
  d <- data.frame(
    weight = c(ctl, trt),
    group = gl(2, 10, 20, labels = c("Ctl", "Trt"))
  )
}
```

```

)
fit <- stan_glm(weight ~ group, data = d, refresh = 0)
yrep <- posterior_predict(fit)
dim(yrep)

log_ratios <- -1 * log_lik(fit)
dim(log_ratios)

r_eff <- relative_eff(exp(-log_ratios), chain_id = rep(1:4, each = 1000))
psis_object <- psis(log_ratios, r_eff = r_eff, cores = 2)

E_loo(yrep, psis_object, type = "mean")
E_loo(yrep, psis_object, type = "var")
E_loo(yrep, psis_object, type = "sd")
E_loo(yrep, psis_object, type = "quantile", probs = 0.5) # median
E_loo(yrep, psis_object, type = "quantile", probs = c(0.1, 0.9))

# We can get more accurate Pareto k diagnostic if we also provide
# the log_ratios argument
E_loo(yrep, psis_object, type = "mean", log_ratios = log_ratios)
}

```

elpd

Generic (expected) log-predictive density

Description

The `elpd()` methods for arrays and matrices can compute the expected log pointwise predictive density for a new dataset or the log pointwise predictive density of the observed data (an overestimate of the elpd).

Usage

```

elpd(x, ...)

## S3 method for class 'array'
elpd(x, ...)

## S3 method for class 'matrix'
elpd(x, ...)

```

Arguments

<code>x</code>	A log-likelihood array or matrix. The Methods (by class) section, below, has detailed descriptions of how to specify the inputs for each method.
<code>...</code>	Currently ignored.

Details

The `elpd()` function is an S3 generic and methods are provided for 3-D pointwise log-likelihood arrays and matrices.

Methods (by class)

- `elpd(array)`: An I by C by N array, where I is the number of MCMC iterations per chain, C is the number of chains, and N is the number of data points.
- `elpd(matrix)`: An S by N matrix, where S is the size of the posterior sample (with all chains merged) and N is the number of data points.

See Also

The vignette *Holdout validation and K-fold cross-validation of Stan programs with the loo package* for demonstrations of using the `elpd()` methods.

Examples

```
# Calculate the lpd of the observed data
LLarr <- example_loglik_array()
elpd(LLarr)
```

example_loglik_array *Objects to use in examples and tests*

Description

Example pointwise log-likelihood objects to use in demonstrations and tests. See the **Value** and **Examples** sections below.

Usage

```
example_loglik_array()
```

```
example_loglik_matrix()
```

Value

`example_loglik_array()` returns a 500 (draws) x 2 (chains) x 32 (observations) pointwise log-likelihood array.

`example_loglik_matrix()` returns the same pointwise log-likelihood values as `example_loglik_array()` but reshaped into a 1000 (draws*chains) x 32 (observations) matrix.

Examples

```
LLarr <- example_loglik_array()
(dim_arr <- dim(LLarr))
LLmat <- example_loglik_matrix()
(dim_mat <- dim(LLmat))

all.equal(dim_mat[1], dim_arr[1] * dim_arr[2])
all.equal(dim_mat[2], dim_arr[3])

all.equal(LLarr[, 1, ], LLmat[1:500, ])
all.equal(LLarr[, 2, ], LLmat[501:1000, ])
```

extract_log_lik

Extract pointwise log-likelihood from a Stan model

Description

Convenience function for extracting the pointwise log-likelihood matrix or array from a stanfit object from the **rstan** package. Note: recent versions of **rstan** now include a loo() method for stanfit objects that handles this internally.

Usage

```
extract_log_lik(stanfit, parameter_name = "log_lik", merge_chains = TRUE)
```

Arguments

stanfit	A stanfit object (rstan package).
parameter_name	A character string naming the parameter (or generated quantity) in the Stan model corresponding to the log-likelihood.
merge_chains	If TRUE (the default), all Markov chains are merged together (i.e., stacked) and a matrix is returned. If FALSE they are kept separate and an array is returned.

Details

Stan does not automatically compute and store the log-likelihood. It is up to the user to incorporate it into the Stan program if it is to be extracted after fitting the model. In a Stan model, the pointwise log likelihood can be coded as a vector in the transformed parameters block (and then summed up in the model block) or it can be coded entirely in the generated quantities block. We recommend using the generated quantities block so that the computations are carried out only once per iteration rather than once per HMC leapfrog step.

For example, the following is the generated quantities block for computing and saving the log-likelihood for a linear regression model with N data points, outcome y, predictor matrix X, coefficients beta, and standard deviation sigma:

```
vector[N] log_lik;
for (n in 1:N) log_lik[n] = normal_lpdf(y[n] | X[n, ] * beta, sigma);
```

Value

If merge_chains=TRUE, an S by N matrix of (post-warmup) extracted draws, where S is the size of the posterior sample and N is the number of data points. If merge_chains=FALSE, an I by C by N array, where $I \times C = S$.

References

Stan Development Team (2017). The Stan C++ Library, Version 2.16.0. <https://mc-stan.org/>
 Stan Development Team (2017). RStan: the R interface to Stan, Version 2.16.1. <https://mc-stan.org/>

find_model_names	<i>Find the model names associated with "loo" objects</i>
------------------	---

Description

Find the model names associated with "loo" objects

Usage

```
find_model_names(x)
```

Arguments

`x` List of "loo" objects.

Value

Character vector of model names the same length as `x`.

gpdfit	<i>Estimate parameters of the Generalized Pareto distribution</i>
--------	---

Description

Given a sample x , Estimate the parameters k and σ of the generalized Pareto distribution (GPD), assuming the location parameter is 0. By default the fit uses a prior for k , which will stabilize estimates for very small sample sizes (and low effective sample sizes in the case of MCMC samples). The weakly informative prior is a Gaussian prior centered at 0.5.

Usage

```
gpdfit(x, wip = TRUE, min_grid_pts = 30, sort_x = TRUE)
```

Arguments

<code>x</code>	A numeric vector. The sample from which to estimate the parameters.
<code>wip</code>	Logical indicating whether to adjust k based on a weakly informative Gaussian prior centered on 0.5. Defaults to TRUE.
<code>min_grid_pts</code>	The minimum number of grid points used in the fitting algorithm. The actual number used is <code>min_grid_pts + floor(sqrt(length(x)))</code> .
<code>sort_x</code>	If TRUE (the default), the first step in the fitting algorithm is to sort the elements of <code>x</code> . If <code>x</code> is already sorted in ascending order then <code>sort_x</code> can be set to FALSE to skip the initial sorting step.

Details

Here the parameter k is the negative of k in Zhang & Stephens (2009).

Value

A named list with components `k` and `sigma`.

References

Zhang, J., and Stephens, M. A. (2009). A new and efficient estimation method for the generalized Pareto distribution. *Technometrics* **51**, 316-325.

See Also

[psis\(\)](#), [pareto-k-diagnostic](#)

`importance_sampling` *A parent class for different importance sampling methods.*

Description

A parent class for different importance sampling methods.

Usage

```
importance_sampling(log_ratios, method, ...)
```

```
## S3 method for class 'array'
importance_sampling(
  log_ratios,
  method,
  ...,
  r_eff = 1,
  cores = getOption("mc.cores", 1)
)
```



```
## S3 method for class 'matrix'
importance_sampling(
  log_ratios,
  method,
  ...,
  r_eff = 1,
  cores = getOption("mc.cores", 1)
)

## Default S3 method:
importance_sampling(log_ratios, method, ..., r_eff = 1)
```

Arguments

log_ratios	An array, matrix, or vector of importance ratios on the log scale (for PSIS-LOO these are <i>negative</i> log-likelihood values). See the Methods (by class) section below for a detailed description of how to specify the inputs for each method.
method	The importance sampling method to use. The following methods are implemented: • "psis": Pareto-Smoothed Importance Sampling (PSIS). Default method. • "tis": Truncated Importance Sampling (TIS) with truncation at $\sqrt{S}$, where S is the number of posterior draws. • "sis": Standard Importance Sampling (SIS).
...	Arguments passed on to the various methods.
r_eff	Vector of relative effective sample size estimates containing one element per observation. The values provided should be the relative effective sample sizes of $1/\exp(\text{log_ratios})$ (i.e., $1/\text{ratios}$). This is related to the relative efficiency of estimating the normalizing term in self-normalizing importance sampling. If <code>r_eff</code> is not provided then the reported PSIS effective sample sizes and Monte Carlo error estimates can be over-optimistic. If the posterior draws are (near) independent then <code>r_eff=1</code> can be used. <code>r_eff</code> has to be a scalar (same value is used for all observations) or a vector with length equal to the number of observations. The default value is 1. See the relative_eff() helper function for computing <code>r_eff</code> .
cores	The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code> . The old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until <code>loo.cores</code> is removed in a future release. As of version 2.0.0 the default is now 1 core if <code>mc.cores</code> is not set , but we recommend using as many (or close to as many) cores as possible. • Note for Windows 10 users: it is strongly recommended to avoid using the <code>.Rprofile</code> file to set <code>mc.cores</code> (using the <code>cores</code> argument or setting <code>mc.cores</code> interactively or in a script is fine).

kfold-generic

Generic function for K-fold cross-validation for developers

Description

For developers of Bayesian modeling packages, **loo** includes a generic function `kfold()` so that methods may be defined for K-fold CV without name conflicts between packages. See, for example, the `kfold()` methods in the **rstanarm** and **brms** packages.

The **Value** section below describes the objects that `kfold()` methods should return in order to be compatible with `loo_compare()` and the **loo** package print methods.

Usage

```
kfold(x, ...)
```

```
is.kfold(x)
```

Arguments

`x` A fitted model object.

`...` Arguments to pass to specific methods.

Value

For developers defining a `kfold()` method for a class "foo", the `kfold.foo()` function should return a list with class `c("kfold", "loo")` with at least the following named elements:

- "estimates": A 1x2 matrix containing the ELPD estimate and its standard error. The matrix must have row name "elpd_kfold" and column names "Estimate" and "SE".
- "pointwise": A Nx1 matrix with column name "elpd_kfold" containing the pointwise contributions for each data point.

It is important for the object to have at least these classes and components so that it is compatible with other functions like `loo_compare()` and `print()` methods.

kfold-helpers

Helper functions for K-fold cross-validation

Description

These functions can be used to generate indexes for use with K-fold cross-validation. See the **Details** section for explanations.

Usage

```
kfold_split_random(K = 10, N = NULL)

kfold_split_stratified(K = 10, x = NULL)

kfold_split_grouped(K = 10, x = NULL)
```

Arguments

K	The number of folds to use.
N	The number of observations in the data.
x	A discrete variable of length N with at least K levels (unique values). Will be coerced to a factor .

Details

`kfold_split_random()` splits the data into K groups of equal size (or roughly equal size).

For a categorical variable `x` `kfold_split_stratified()` splits the observations into K groups ensuring that relative category frequencies are approximately preserved.

For a grouping variable `x`, `kfold_split_grouped()` places all observations in `x` from the same group/level together in the same fold. The selection of which groups/levels go into which fold (relevant when there are more groups than folds) is randomized.

Value

An integer vector of length N where each element is an index in 1:K.

Examples

```
ids <- kfold_split_random(K = 5, N = 20)
print(ids)
table(ids)

x <- sample(c(0, 1), size = 200, replace = TRUE, prob = c(0.05, 0.95))
table(x)
ids <- kfold_split_stratified(K = 5, x = x)
print(ids)
table(ids, x)

grp <- gl(n = 50, k = 15, labels = state.name)
length(grp)
head(table(grp))

ids_10 <- kfold_split_grouped(K = 10, x = grp)
(tab_10 <- table(grp, ids_10))
colSums(tab_10)

ids_9 <- kfold_split_grouped(K = 9, x = grp)
```

```
(tab_9 <- table(grp, ids_9))
colSums(tab_9)
```

loo

Efficient approximate leave-one-out cross-validation (LOO)

Description

The `loo()` methods for arrays, matrices, and functions compute PSIS-LOO CV, efficient approximate leave-one-out (LOO) cross-validation for Bayesian models using Pareto smoothed importance sampling ([PSIS](#)). This is an implementation of the methods described in Vehtari, Gelman, and Gabry (2017) and Vehtari, Simpson, Gelman, Yao, and Gabry (2024).

The `loo_i()` function enables testing log-likelihood functions for use with the `loo.function()` method.

Usage

```
loo(x, ...)

## S3 method for class 'array'
loo(
  x,
  ...,
  r_eff = 1,
  save_psis = FALSE,
  cores = getOption("mc.cores", 1),
  is_method = c("psis", "tis", "sis")
)

## S3 method for class 'matrix'
loo(
  x,
  ...,
  r_eff = 1,
  save_psis = FALSE,
  cores = getOption("mc.cores", 1),
  is_method = c("psis", "tis", "sis")
)

## S3 method for class '`function`'
loo(
  x,
  ...,
  data = NULL,
  draws = NULL,
  r_eff = 1,
```

```

    save_psis = FALSE,
    cores = getOption("mc.cores", 1),
    is_method = c("psis", "tis", "sis")
  )

loo_i(i, llfun, ..., data = NULL, draws = NULL, r_eff = 1, is_method = "psis")

is.loo(x)

is.psis_loo(x)

```

Arguments

<code>x</code>	A log-likelihood array, matrix, or function. The Methods (by class) section, below, has detailed descriptions of how to specify the inputs for each method.
<code>r_eff</code>	Vector of relative effective sample size estimates for the likelihood ($\exp(\log_lik)$) of each observation. This is related to the relative efficiency of estimating the normalizing term in self-normalized importance sampling when using posterior draws obtained with MCMC. If MCMC draws are used and <code>r_eff</code> is not provided then the reported PSIS effective sample sizes and Monte Carlo error estimates can be over-optimistic. If the posterior draws are (near) independent then <code>r_eff=1</code> can be used. <code>r_eff</code> has to be a scalar (same value is used for all observations) or a vector with length equal to the number of observations. The default value is 1. See the relative_eff() helper functions for help computing <code>r_eff</code> .
<code>save_psis</code>	Should the <code>psis</code> object created internally by <code>loo()</code> be saved in the returned object? The <code>loo()</code> function calls psis() internally but by default discards the (potentially large) <code>psis</code> object after using it to compute the LOO-CV summaries. Setting <code>save_psis=TRUE</code> will add a <code>psis_object</code> component to the list returned by <code>loo</code> . This is useful if you plan to use the E_loo() function to compute weighted expectations after running <code>loo</code> . Several functions in the bayesplot package also accept <code>psis</code> objects.
<code>cores</code>	The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code> . The old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until <code>loo.cores</code> is removed in a future release. As of version 2.0.0 the default is now 1 core if <code>mc.cores</code> is not set , but we recommend using as many (or close to as many) cores as possible. • Note for Windows 10 users: it is strongly recommended to avoid using the <code>.Rprofile</code> file to set <code>mc.cores</code> (using the <code>cores</code> argument or setting <code>mc.cores</code> interactively or in a script is fine).
<code>is_method</code>	The importance sampling method to use. The following methods are implemented: • "psis": Pareto-Smoothed Importance Sampling (PSIS). Default method. • "tis": Truncated Importance Sampling (TIS) with truncation at $\sqrt{t(S)}$, where S is the number of posterior draws. • "sis": Standard Importance Sampling (SIS).

data, draws, ...	For the <code>loo.function()</code> method and the <code>loo_i()</code> function, these are the data, posterior draws, and other arguments to pass to the log-likelihood function. See the Methods (by class) section below for details on how to specify these arguments.
i	For <code>loo_i()</code> , an integer in $1:N$.
llfun	For <code>loo_i()</code> , the same as <code>x</code> for the <code>loo.function()</code> method. A log-likelihood function as described in the Methods (by class) section.

Details

The `loo()` function is an S3 generic and methods are provided for 3-D pointwise log-likelihood arrays, pointwise log-likelihood matrices, and log-likelihood functions. The array and matrix methods are the most convenient, but for models fit to very large datasets the `loo.function()` method is more memory efficient and may be preferable.

Value

The `loo()` methods return a named list with class `c("psis_loo", "loo")` and components:

- estimates** A matrix with two columns (Estimate, SE) and three rows (`elpd_loo`, `p_loo`, `looic`). This contains point estimates and standard errors of the expected log pointwise predictive density ([elpd_loo](#)), the effective number of parameters ([p_loo](#)) and the LOO information criterion `looic` (which is just $-2 * \text{elpd_loo}$, i.e., converted to deviance scale).
- pointwise** A matrix with five columns (and number of rows equal to the number of observations) containing the pointwise contributions of the measures (`elpd_loo`, `mcse_elpd_loo`, `p_loo`, `looic`, `influence_pareto_k`). In addition to the three measures in `estimates`, we also report pointwise values of the Monte Carlo standard error of [elpd_loo](#) ([mcse_elpd_loo](#)), and statistics describing the influence of each observation on the posterior distribution (`influence_pareto_k`). These are the estimates of the shape parameter k of the generalized Pareto fit to the importance ratios for each leave-one-out distribution (see the [pareto-k-diagnostic](#) page for details).
- diagnostics** A named list containing two vectors:
 - `pareto_k`: Importance sampling reliability diagnostics. By default, these are equal to the `influence_pareto_k` in `pointwise`. Some algorithms can improve importance sampling reliability and modify these diagnostics. See the [pareto-k-diagnostic](#) page for details.
 - `n_eff`: PSIS effective sample size estimates.
- psis_object** This component will be NULL unless the `save_psis` argument is set to TRUE when calling `loo()`. In that case `psis_object` will be the object of class "psis" that is created when the `loo()` function calls [psis\(\)](#) internally to do the PSIS procedure.

The `loo_i()` function returns a named list with components `pointwise` and `diagnostics`. These components have the same structure as the `pointwise` and `diagnostics` components of the object returned by `loo()` except they contain results for only a single observation.

Methods (by class)

- `loo(array)`: An I by C by N array, where I is the number of MCMC iterations per chain, C is the number of chains, and N is the number of data points.

- `loo(matrix)`: An S by N matrix, where S is the size of the posterior sample (with all chains merged) and N is the number of data points.
- `loo(`function`)`: A function `f()` that takes arguments `data_i` and `draws` and returns a vector containing the log-likelihood for a single observation i evaluated at each posterior draw. The function should be written such that, for each observation i in $1:N$, evaluating

```
f(data_i = data[i,, drop=FALSE], draws = draws)
```

results in a vector of length S (size of posterior sample). The log-likelihood function can also have additional arguments but `data_i` and `draws` are required.

If using the function method then the arguments `data` and `draws` must also be specified in the call to `loo()`:

- `data`: A data frame or matrix containing the data (e.g. observed outcome and predictors) needed to compute the pointwise log-likelihood. For each observation i , the i th row of `data` will be passed to the `data_i` argument of the log-likelihood function.
- `draws`: An object containing the posterior draws for any parameters needed to compute the pointwise log-likelihood. Unlike `data`, which is indexed by observation, for each observation the entire object `draws` will be passed to the `draws` argument of the log-likelihood function.
- The `...` can be used if your log-likelihood function takes additional arguments. These arguments are used like the `draws` argument in that they are recycled for each observation.

Defining `loo()` methods in a package

Package developers can define `loo()` methods for fitted models objects. See the example `loo.stanfit()` method in the **Examples** section below for an example of defining a method that calls `loo.array()`. The `loo.stanreg()` method in the **rstanarm** package is an example of defining a method that calls `loo.function()`.

References

- Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).
- Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)

See Also

- The **loo** package [vignettes](#) for demonstrations.
- The [FAQ page](#) on the **loo** website for answers to frequently asked questions.
- [psis\(\)](#) for the underlying Pareto Smoothed Importance Sampling (PSIS) procedure used in the LOO-CV approximation.
- [pareto-k-diagnostic](#) for convenience functions for looking at diagnostics.
- [loo_compare\(\)](#) for model comparison.

Examples

```

### Array and matrix methods (using example objects included with loo package)
# Array method
LLarr <- example_loglik_array()
rel_n_eff <- relative_eff(exp(LLarr))
loo(LLarr, r_eff = rel_n_eff, cores = 2)

# Matrix method
LLmat <- example_loglik_matrix()
rel_n_eff <- relative_eff(exp(LLmat), chain_id = rep(1:2, each = 500))
loo(LLmat, r_eff = rel_n_eff, cores = 2)

### Using log-likelihood function instead of array or matrix
set.seed(124)

# Simulate data and draw from posterior
N <- 50; K <- 10; S <- 100; a0 <- 3; b0 <- 2
p <- rbeta(1, a0, b0)
y <- rbinom(N, size = K, prob = p)
a <- a0 + sum(y); b <- b0 + N * K - sum(y)
fake_posterior <- as.matrix(rbeta(S, a, b))
dim(fake_posterior) # S x 1
fake_data <- data.frame(y, K)
dim(fake_data) # N x 2

llfun <- function(data_i, draws) {
  # each time called internally within loo the arguments will be equal to:
  # data_i: ith row of fake_data (fake_data[i,, drop=FALSE])
  # draws: entire fake_posterior matrix
  dbinom(data_i$y, size = data_i$K, prob = draws, log = TRUE)
}

# Use the loo_i function to check that llfun works on a single observation
# before running on all obs. For example, using the 3rd obs in the data:
loo_3 <- loo_i(i = 3, llfun = llfun, data = fake_data, draws = fake_posterior)
print(loo_3$pointwise[, "elpd_loo"])

# Use loo.function method (default r_eff=1 is used as this posterior not obtained via MCMC)
loo_with_fn <- loo(llfun, draws = fake_posterior, data = fake_data)

# If we look at the elpd_loo contribution from the 3rd obs it should be the
# same as what we got above with the loo_i function and i=3:
print(loo_with_fn$pointwise[3, "elpd_loo"])
print(loo_3$pointwise[, "elpd_loo"])

# Check that the loo.matrix method gives same answer as loo.function method
log_lik_matrix <- sapply(1:N, function(i) {
  llfun(data_i = fake_data[i,, drop=FALSE], draws = fake_posterior)
})
loo_with_mat <- loo(log_lik_matrix)
all.equal(loo_with_mat$estimates, loo_with_fn$estimates) # should be TRUE!

```



```
## Not run:
### For package developers: defining loo methods

# An example of a possible loo method for 'stanfit' objects (rstan package).
# A similar method is included in the rstan package.
# In order for users to be able to call loo(stanfit) instead of
# loo.stanfit(stanfit) the NAMESPACE needs to be handled appropriately
# (roxygen2 and devtools packages are good for that).
#
loo.stanfit <-
  function(x,
           pars = "log_lik",
           ...,
           save_psis = FALSE,
           cores = getOption("mc.cores", 1)) {
    stopifnot(length(pars) == 1L)
    LLarray <- loo::extract_log_lik(stanfit = x,
                                   parameter_name = pars,
                                   merge_chains = FALSE)
    r_eff <- loo::relative_eff(x = exp(LLarray), cores = cores)
    loo::loo.array(LLarray,
                   r_eff = r_eff,
                   cores = cores,
                   save_psis = save_psis)
  }

## End(Not run)
```

 loo-datasets

Datasets for loo examples and vignettes

Description

Small datasets for use in **loo** examples and vignettes. The Kline and milk datasets are also included in the **rethinking** package (McElreath, 2016a), but we include them here as **rethinking** is not on CRAN.

Details

Currently the data sets included are:

- Kline: Small dataset from Kline and Boyd (2010) on tool complexity and demography in Oceanic islands societies. This data is discussed in detail in McElreath (2016a,2016b). ([Link to variable descriptions](#))
- milk: Small dataset from Hinde and Milligan (2011) on primate milk composition. This data is discussed in detail in McElreath (2016a,2016b). ([Link to variable descriptions](#))
- voice: Voice rehabilitation data from Tsanas et al. (2014).

References

- Hinde and Milligan. 2011. *Evolutionary Anthropology* 20:9-23.
- Kline, M.A. and R. Boyd. 2010. *Proc R Soc B* 277:2559-2564.
- McElreath, R. (2016a). rethinking: Statistical Rethinking book package. R package version 1.59.
- McElreath, R. (2016b). *Statistical rethinking: A Bayesian course with examples in R and Stan*. Chapman & Hall/CRC.
- A. Tsanas, M.A. Little, C. Fox, L.O. Ramig: Objective automatic assessment of rehabilitative speech treatment in Parkinson's disease, *IEEE Transactions on Neural Systems and Rehabilitation Engineering*, Vol. 22, pp. 181-190, January 2014

Examples

```
str(Kline)
str(milk)
```

 loo-glossary

LOO package glossary

Description

This page provides definitions of key terms. Also see the [FAQ page](#) on the **loo** website for answers to frequently asked questions.

Note: VGG2017 refers to Vehtari, Gelman, and Gabry (2017). See **References**, below.

ELPD and `elpd_loo`

The ELPD is the theoretical expected log pointwise predictive density for a new dataset (Eq 1 in VGG2017), which can be estimated, e.g., using cross-validation. `elpd_loo` is the Bayesian LOO estimate of the expected log pointwise predictive density (Eq 4 in VGG2017) and is a sum of N individual pointwise log predictive densities. Probability densities can be smaller or larger than 1, and thus log predictive densities can be negative or positive. For simplicity the ELPD acronym is used also for expected log pointwise predictive probabilities for discrete models. Probabilities are always equal or less than 1, and thus log predictive probabilities are 0 or negative.

Standard error of `elpd_loo`

As `elpd_loo` is defined as the sum of N independent components (Eq 4 in VGG2017), we can compute the standard error by using the standard deviation of the N components and multiplying by $\sqrt{N}$ (Eq 23 in VGG2017). This standard error is a coarse description of our uncertainty about the predictive performance for unknown future data. When N is small or there is severe model misspecification, the current SE estimate is overoptimistic and the actual SE can even be twice as large. Even for moderate N , when the SE estimate is an accurate estimate for the scale, it ignores the skewness. When making model comparisons, the SE of the component-wise (pairwise) differences should be used instead (see the `se_diff` section below and Eq 24 in VGG2017). Sivula et al. (2022) discuss the conditions when the normal approximation used for SE and `se_diff` is good.

Monte Carlo SE of elpd_{loo}

The Monte Carlo standard error is the estimate for the computational accuracy of MCMC and importance sampling used to compute elpd_{loo} . Usually this is negligible compared to the standard describing the uncertainty due to finite number of observations (Eq 23 in VGG2017).

p_{loo} (effective number of parameters)

p_{loo} is the difference between elpd_{loo} and the non-cross-validated log posterior predictive density. It describes how much more difficult it is to predict future data than the observed data. Asymptotically under certain regularity conditions, p_{loo} can be interpreted as the *effective number of parameters*. In well behaving cases $p_{\text{loo}} < N$ and $p_{\text{loo}} < p$, where p is the total number of parameters in the model. $p_{\text{loo}} > N$ or $p_{\text{loo}} > p$ indicates that the model has very weak predictive capability and may indicate a severe model misspecification. See below for more on interpreting p_{loo} when there are warnings about high Pareto k diagnostic values.

Pareto k estimates

The Pareto k estimate is a diagnostic for Pareto smoothed importance sampling (PSIS), which is used to compute components of elpd_{loo} . In importance-sampling LOO the full posterior distribution is used as the proposal distribution. The Pareto k diagnostic estimates how far an individual leave-one-out distribution is from the full distribution. If leaving out an observation changes the posterior too much then importance sampling is not able to give a reliable estimate. Pareto smoothing stabilizes importance sampling and guarantees a finite variance estimate at the cost of some bias.

The diagnostic threshold for Pareto k depends on sample size S (sample size dependent threshold was introduced by Vehtari et al., 2024, and before that fixed thresholds of 0.5 and 0.7 were recommended). For simplicity, loo package uses the nominal sample size S when computing the sample size specific threshold. This provides an optimistic threshold if the effective sample size is less than 2200, but even then if $\text{ESS}/S > 1/2$ the difference is usually negligible. Thinning of MCMC draws can be used to improve the ratio ESS/S .

- If $k < \min(1 - 1/\log_{10}(S), 0.7)$, where S is the sample size, the PSIS estimate and the corresponding Monte Carlo standard error estimate are reliable.
- If $1 - 1/\log_{10}(S) \leq k < 0.7$, the PSIS estimate and the corresponding Monte Carlo standard error estimate are not reliable, but increasing the (effective) sample size S above 2200 may help (this will increase the sample size specific threshold $(1 - 1/\log_{10}(2200)) > 0.7$ and then the bias specific threshold 0.7 dominates).
- If $0.7 \leq k < 1$, the PSIS estimate and the corresponding Monte Carlo standard error have large bias and are not reliable. Increasing the sample size may reduce the variability in the k estimate, which may also result in a lower k estimate.
- If $k \geq 1$, the target distribution is estimated to have non-finite mean. The PSIS estimate and the corresponding Monte Carlo standard error are not well defined. Increasing the sample size may reduce the variability in k estimate, which may also result in a lower k estimate.

Pareto k is also useful as a measure of influence of an observation. Highly influential observations have high k values. Very high k values often indicate model misspecification, outliers or mistakes in data processing. See Section 6 of Gabry et al. (2019) for an example.

Interpreting p_{loo} when Pareto k is large: If $k > 0.7$ then we can also look at the p_{loo} estimate for some additional information about the problem:

- If $p_{loo} \ll p$ (the total number of parameters in the model), then the model is likely to be misspecified. Posterior predictive checks (PPCs) are then likely to also detect the problem. Try using an overdispersed model, or add more structural information (nonlinearity, mixture model, etc.).
- If $p_{loo} < p$ and the number of parameters p is relatively large compared to the number of observations (e.g., $p > N/5$), it is likely that the model is so flexible or the population prior so weak that it's difficult to predict the left out observation (even for the true model). This happens, for example, in the simulated 8 schools (in VGG2017), random effect models with a few observations per random effect, and Gaussian processes and spatial models with short correlation lengths.
- If $p_{loo} > p$, then the model is likely to be badly misspecified. If the number of parameters $p \ll N$, then PPCs are also likely to detect the problem. See the case study at <https://avehtari.github.io/modelselection/roaches.html> for an example. If p is relatively large compared to the number of observations, say $p > N/5$ (more accurately we should count number of observations influencing each parameter as in hierarchical models some groups may have few observations and other groups many), it is possible that PPCs won't detect the problem.

elpd_diff

$elpd_diff$ is the difference in $elpd_{loo}$ for two models. If more than two models are compared, the difference is computed relative to the model with highest $elpd_{loo}$.

se_diff

The standard error of component-wise differences of $elpd_{loo}$ (Eq 24 in VGG2017) between two models. This SE is *smaller* than the SE for individual models due to correlation (i.e., if some observations are easier and some more difficult to predict for all models).

p_{worse} (probability of worse predictive performance)

p_{worse} is the estimated probability that a model has worse predictive performance than the best-ranked model in the comparison, based on the normal approximation to the uncertainty in $elpd_diff$. It is computed as

```
p_worse = pnorm(0, elpd_diff, se_diff).
```

The best-ranked model (the first row in the `loo_compare()` output, where $elpd_diff = 0$) always receives NA, since the comparison is defined relative to that model.

Because models are ordered by $elpd_{loo}$ before computing p_{worse} , all reported values are at least 0.5 by construction. A value close to 0.5 indicates that the models are nearly indistinguishable in predictive performance and that the ranking could easily be reversed with different data. A value close to 1 indicates that the lower-ranked model is almost certainly worse. p_{worse} inherits all the limitations of se_diff and the normal approximation on which it is based. In particular, when se_diff is underestimated, p_{worse} will be estimated too close to 1, making a model appear more clearly worse than the data actually support. Conversely, when $elpd_diff$ is biased due to

an unreliable LOO approximation, `p_worse` can point in the wrong direction entirely. When any of these conditions are present, `diag_diff` or `diag_elpd` will be flagged in the `loo_compare()` output. For further guidance, see the sections below and the case study on [Uncertainty in Bayesian LOO-CV Model Comparison](#).

`diag_diff` (pairwise comparison diagnostics)

`diag_diff` is a diagnostic column in the `loo_compare()` output for each model comparison against the current reference model. It flags conditions under which the normal approximation behind `se_diff` and `p_worse` is likely to be poorly calibrated. The column contains a short label when a condition is detected, and is empty otherwise.

The column `diag_diff` currently flags two problems:

`N < 100`:

When the number of observations is small, we may assume `se_diff` to be underestimated. As a rough heuristic one can multiply `se_diff` by 2 to make a more conservative estimate.

`|elpd_diff| < 4`:

When `|elpd_diff|` is below 4, the models have very similar predictive performance. In this setting, Sivula et al. (2025) show that skewness in the error distribution can make the normal approximation for `se_diff` and `p_worse` miscalibrated, even for large `N`. In practice, this usually supports treating the models as predictively similar.

Relation between `N < 100` and `|elpd_diff| < 4`:

The conditions flagged by `diag_diff` are not independent: they tend to co-occur, and when they do, some flags carry more information than others. `loo_compare()` therefore follows a priority hierarchy and shows only the most critical flag in the table output.

The hierarchy is as follows:

- **`N < 100` takes highest priority.** A small sample size undermines the reliability of `se_diff` by underestimating uncertainty. Because of this, even if `|elpd_diff| < 4` is also true for a comparison, the table will only show `N < 100`. The small sample size renders the `|elpd_diff| < 4` diagnostic less meaningful.
- **`|elpd_diff| < 4` takes second priority.** When `N ≥ 100` and the difference is small, the normal approximation is miscalibrated due to the skewness of the error distribution (Sivula et al., 2025). In this situation, `se_diff` exists and is not heavily biased in scale, but the shape of the approximation is wrong, making `p_worse` unreliable.

For further guidance, see the case study on [Uncertainty in Bayesian LOO-CV Model Comparison](#).

`diag_elpd`

`diag_elpd` is a diagnostic column in the `loo_compare()` output that flags when the PSIS-LOO approximation for an individual model is unreliable. Unlike `diag_diff`, which concerns the *comparison* between models, `diag_elpd` concerns the quality of the `elpd_loo` estimate for each model individually. It contains a short text label when a problem is detected, and is empty otherwise.

`K k_psis > t` (K observations with Pareto-k values > t):

This label indicates that `K` observations for this model have Pareto-k values above the PSIS reliability threshold `t` used by `loo` for that fit. The threshold is sample-size dependent, and in many

practical cases close to 0.7. When this flag appears, the PSIS approximation can be unreliable for those observations, and the resulting `elpd_loo` may be biased. Because `elpd_diff` is a direct difference of two models' `elpd_loo` values, bias in either model's estimate propagates directly into `elpd_diff` and `p_worse`. This is qualitatively different from the calibration issues flagged by `diag_diff`: here the estimate itself may be wrong, not just uncertain.

See for further information on Pareto-k values the "Pareto k estimates" section.

References

- Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*. 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).
- Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)
- Sivula, T, Magnusson, M., Matamoros A. A., and Vehtari, A. (2025). Uncertainty in Bayesian leave-one-out cross-validation based model comparison. *Bayesian Analysis*. doi:10.1214/25BA1569.
- Gabry, J. , Simpson, D. , Vehtari, A. , Betancourt, M. and Gelman, A. (2019), Visualization in Bayesian workflow. *J. R. Stat. Soc. A*, 182: 389-402. doi:10.1111/rssa.12378 ([journal version](#), [preprint arXiv:1709.01449](#), [code on GitHub](#))

loo_approximate_posterior

Efficient approximate leave-one-out cross-validation (LOO) for posterior approximations

Description

Efficient approximate leave-one-out cross-validation (LOO) for posterior approximations

Usage

```
loo_approximate_posterior(x, log_p, log_g, ...)
```

```
## S3 method for class 'array'
loo_approximate_posterior(
  x,
  log_p,
  log_g,
  ...,
  save_psis = FALSE,
  cores = getOption("mc.cores", 1)
)
```

```
## S3 method for class 'matrix'
loo_approximate_posterior(
  x,
```

```

    log_p,
    log_g,
    ...,
    save_psis = FALSE,
    cores = getOption("mc.cores", 1)
)

## S3 method for class ``function``
loo_approximate_posterior(
  x,
  ...,
  data = NULL,
  draws = NULL,
  log_p = NULL,
  log_g = NULL,
  save_psis = FALSE,
  cores = getOption("mc.cores", 1)
)

```

Arguments

<code>x</code>	A log-likelihood array, matrix, or function. The Methods (by class) section, below, has detailed descriptions of how to specify the inputs for each method.
<code>log_p</code>	The log-posterior (target) evaluated at S samples from the proposal distribution (g). A vector of length S .
<code>log_g</code>	The log-density (proposal) evaluated at S samples from the proposal distribution (g). A vector of length S .
<code>save_psis</code>	Should the "psis" object created internally by <code>loo_approximate_posterior()</code> be saved in the returned object? See <code>loo()</code> for details.
<code>cores</code>	The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code> . The old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until <code>loo.cores</code> is removed in a future release. As of version 2.0.0 the default is now 1 core if <code>mc.cores</code> is not set , but we recommend using as many (or close to as many) cores as possible. Note for Windows 10 users: it is strongly recommended to avoid using the <code>.Rprofile</code> file to set <code>mc.cores</code> (using the <code>cores</code> argument or setting <code>mc.cores</code> interactively or in a script is fine).
<code>data, draws, ...</code>	For the <code>loo_approximate_posterior.function()</code> method, these are the data, posterior draws, and other arguments to pass to the log-likelihood function. See the Methods (by class) section below for details on how to specify these arguments.

Details

The `loo_approximate_posterior()` function is an S3 generic and methods are provided for 3-D pointwise log-likelihood arrays, pointwise log-likelihood matrices, and log-likelihood functions.

The implementation works for posterior approximations where it is possible to compute the log density for the posterior approximation.

Value

The `loo_approximate_posterior()` methods return a named list with class `c("psis_loo_ap", "psis_loo", "loo")`. It has the same structure as the objects returned by `loo()` but with the additional slot:

`posterior_approximation` A list with two vectors, `log_p` and `log_g` of the same length containing the posterior density and the approximation density for the individual draws.

Methods (by class)

- `loo_approximate_posterior(array)`: An I by C by N array, where I is the number of MCMC iterations per chain, C is the number of chains, and N is the number of data points.
- `loo_approximate_posterior(matrix)`: An S by N matrix, where S is the size of the posterior sample (with all chains merged) and N is the number of data points.
- `loo_approximate_posterior(`function`)`: A function `f()` that takes arguments `data_i` and `draws` and returns a vector containing the log-likelihood for a single observation i evaluated at each posterior draw. The function should be written such that, for each observation i in $1:N$, evaluating

```
f(data_i = data[i,, drop=FALSE], draws = draws)
```

results in a vector of length S (size of posterior sample). The log-likelihood function can also have additional arguments but `data_i` and `draws` are required.

If using the function method then the arguments `data` and `draws` must also be specified in the call to `loo()`:

- `data`: A data frame or matrix containing the data (e.g. observed outcome and predictors) needed to compute the pointwise log-likelihood. For each observation i , the i th row of `data` will be passed to the `data_i` argument of the log-likelihood function.
- `draws`: An object containing the posterior draws for any parameters needed to compute the pointwise log-likelihood. Unlike `data`, which is indexed by observation, for each observation the entire object `draws` will be passed to the `draws` argument of the log-likelihood function.
- The `...` can be used if your log-likelihood function takes additional arguments. These arguments are used like the `draws` argument in that they are recycled for each observation.

References

Magnusson, M., Riis Andersen, M., Jonasson, J. and Vehtari, A. (2019). Leave-One-Out Cross-Validation for Large Data. In *Thirty-sixth International Conference on Machine Learning*, PMLR 97:4244-4253.

Magnusson, M., Riis Andersen, M., Jonasson, J. and Vehtari, A. (2020). Leave-One-Out Cross-Validation for Model Comparison in Large Data. In *Proceedings of the 23rd International Conference on Artificial Intelligence and Statistics (AISTATS)*, PMLR 108:341-351.

See Also

[loo\(\)](#), [psis\(\)](#), [loo_compare\(\)](#)

loo_compare	<i>Model comparison</i>
-------------	-------------------------

Description

Compare fitted models based on [ELPD](#).

Usage

```
loo_compare(x, ...)
```

Default S3 method:

```
loo_compare(x, ...)
```

S3 method for class 'compare.loo'

```
print(x, ..., digits = 1, p_worse = TRUE, simplify = TRUE)
```

S3 method for class 'compare.loo_ss'

```
print(x, ..., digits = 1)
```

Arguments

x	An object of class "loo" or a list of such objects. If a list is used then the list names will be used as the model names in the output. See Examples .
...	Additional objects of class "loo", if not passed in as a single list.
digits	For the print method only, the number of digits to use when printing.
p_worse	For the print method only, should we include the normal approximation based probability of each model having worse performance than the best model? The default is TRUE.
simplify	For the print method only, should the output be simplified to only include the model names, ELPD differences, and (when p_worse = TRUE) diagnostic columns? The default is TRUE. Set to FALSE to also print the available estimate columns (pointwise ELPD, LOOIC/WAIC, and their standard errors).

Details

When comparing two fitted models, we can estimate the difference in their expected predictive accuracy by the difference in [elpd_loo](#) or [elpd_waic](#) (or multiplied by -2 , if desired, to be on the deviance scale).

elpd_diff and se_diff:

When using `loo_compare()`, the returned data frame will have one row per model and several columns of estimates. The values of `elpd_diff` and `se_diff` are computed by making pairwise comparisons between each model and the model with the largest ELPD (the model listed first). Therefore, the first `elpd_diff` value will always be 0 (i.e., the difference between the preferred model and itself) and the rest of the values will be negative.

To compute the standard error of the difference in ELPD — which should not be expected to equal the difference of the standard errors — we use a paired estimate to take advantage of the fact that the same set of N data points was used to fit both models. These calculations should be most useful when N is large, because then non-normality of the distribution is not such an issue when estimating the uncertainty in these sums. These standard errors, for all their flaws, should give a better sense of uncertainty than what is obtained using the current standard approach of comparing differences of deviances to a Chi-squared distribution, a practice derived for Gaussian linear models or asymptotically, and which only applies to nested models in any case.

p_worse, diag_diff, and diag_elpd:

The values in the `p_worse` column show the probability of each model having worse ELPD than the best model. These probabilities are computed with a normal approximation using the values from `elpd_diff` and `se_diff`. Sivula et al. (2025) present the conditions when the normal approximation used for SE and `se_diff` is good, and the column `diag_diff` contains possible diagnostic messages:

- $N < 100$ (small data)
- $|\text{elpd_diff}| < 4$ (models make similar predictions)

If either of these diagnostic messages is shown, the error distribution is skewed or thick tailed and the normal approximation based on `elpd_diff` and `se_diff` is not well calibrated. In that case, the probabilities `p_worse` are likely to be too large. However, `elpd_diff` and `se_diff` will still be indicative of the differences and uncertainties (for example, if $|\text{elpd_diff}|$ is many times larger than `se_diff` the difference is quite certain). In addition, if the model is not well specified and there are outliers, the error distribution can also be skewed or thick tailed and the normal approximation is not well calibrated. Possible model misspecification and outliers can be diagnosed with usual predictive checking methods.

The column `diag_elpd` shows the PSIS-LOO Pareto k diagnostic for the pointwise ELPD computations for each model. If $k_{\text{psis}} > 0.7$ is shown, where K is the number of high Pareto k values in the PSIS computation, then there may be significant bias in `elpd_diff` favoring models with a large number of high Pareto k values.

Warnings for many model comparisons:

If more than 11 models are compared, we internally recompute the model differences using the median model by ELPD as the baseline model. We then estimate whether the differences in predictive performance are potentially due to chance as described by McLatchie and Vehtari (2023). This will flag a warning if it is deemed that there is a risk of over-fitting due to the selection process. In that case users are recommended to avoid model selection based on LOO-CV, and instead to favor model averaging/stacking or projection predictive inference.

Value

A data frame with class `"compare.loo"` that has its own print method. See the **Details** and **Examples** sections.

References

- Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*. 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).
- Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)
- Sivula, T, Magnusson, M., Matamoros A. A., and Vehtari, A. (2025). Uncertainty in Bayesian leave-one-out cross-validation based model comparison. *Bayesian Analysis*. doi:10.1214/25BA1569
- McLachrie, Y., and Vehtari, A. (2024). Efficient estimation and correction of selection-induced bias with order statistics. *Statistics and Computing*. 34(132). doi:10.1007/s11222024104424

See Also

- The [FAQ page](#) on the **loo** website for answers to frequently asked questions.

Examples

```
# very artificial example, just for demonstration!
LL <- example_loglik_array()
loo1 <- loo(LL)      # should be worst model when compared
loo2 <- loo(LL + 1)  # should be second best model when compared
loo3 <- loo(LL + 2)  # should be best model when compared

comp <- loo_compare(loo1, loo2, loo3)
print(comp, digits = 2)
print(comp, simplify = FALSE) # full table

# can use a list of objects with custom names
# the names will be used in the output
loo_compare(list("apple" = loo1, "banana" = loo2, "cherry" = loo3))

## Not run:
# works for waic (and kfold) too
loo_compare(waic(LL), waic(LL - 10))

## End(Not run)
```

loo_model_weights

Model averaging/weighting via stacking or pseudo-BMA weighting

Description

Model averaging via stacking of predictive distributions, pseudo-BMA weighting or pseudo-BMA+ weighting with the Bayesian bootstrap. See Yao et al. (2018), Vehtari, Gelman, and Gabry (2017), and Vehtari, Simpson, Gelman, Yao, and Gabry (2024) for background.

Usage

```
loo_model_weights(x, ...)

## Default S3 method:
loo_model_weights(
  x,
  ...,
  method = c("stacking", "pseudobma"),
  optim_method = "BFGS",
  optim_control = list(),
  BB = TRUE,
  BB_n = 1000,
  alpha = 1,
  r_eff_list = NULL,
  cores = getOption("mc.cores", 1)
)

stacking_weights(lpd_point, optim_method = "BFGS", optim_control = list())

pseudobma_weights(lpd_point, BB = TRUE, BB_n = 1000, alpha = 1)
```

Arguments

x	A list of "psis_loo" objects (objects returned by <code>loo()</code>) or pointwise log-likelihood matrices, one for each model. If the list elements are named the names will be used to label the models in the results. Each matrix/object should have dimensions S by N , where S is the size of the posterior sample (with all chains merged) and N is the number of data points. If <code>x</code> is a list of log-likelihood matrices then <code>loo()</code> is called internally on each matrix. Currently the <code>loo_model_weights()</code> function is not implemented to be used with results from K-fold CV, but you can still obtain weights using K-fold CV results by calling the <code>stacking_weights()</code> or <code>pseudobma_weights()</code> function directly.
...	Unused, except for the generic to pass arguments to individual methods.
method	Either "stacking" (the default) or "pseudobma", indicating which method to use for obtaining the weights. "stacking" refers to stacking of predictive distributions and "pseudobma" refers to pseudo-BMA+ weighting (or plain pseudo-BMA weighting if argument <code>BB</code> is FALSE).
optim_method	If <code>method="stacking"</code> , a string passed to the <code>method</code> argument of <code>stats::constrOptim()</code> to specify the optimization algorithm. The default is <code>optim_method="BFGS"</code> , but other options are available (see <code>stats::optim()</code>).
optim_control	If <code>method="stacking"</code> , a list of control parameters for optimization passed to the <code>control</code> argument of <code>stats::constrOptim()</code> .
BB	Logical used when <code>method="pseudobma"</code> . If TRUE (the default), the Bayesian bootstrap will be used to adjust the pseudo-BMA weighting, which is called pseudo-BMA+ weighting. It helps regularize the weight away from 0 and 1, so as to reduce the variance.

BB_n	For pseudo-BMA+ weighting only, the number of samples to use for the Bayesian bootstrap. The default is BB_n=1000.
alpha	Positive scalar shape parameter in the Dirichlet distribution used for the Bayesian bootstrap. The default is alpha=1, which corresponds to a uniform distribution on the simplex space.
r_eff_list	Optionally, a list of relative effective sample size estimates for the likelihood ($\exp(\log_lik)$) of each observation in each model. See psis() and relative_eff() helper function for computing r_eff. If x is a list of "psis_loo" objects then r_eff_list is ignored.
cores	The number of cores to use for parallelization. This defaults to the option mc.cores which can be set for an entire R session by options(mc.cores = NUMBER). The old option loo.cores is now deprecated but will be given precedence over mc.cores until loo.cores is removed in a future release. As of version 2.0.0 the default is now 1 core if mc.cores is not set , but we recommend using as many (or close to as many) cores as possible. • Note for Windows 10 users: it is strongly recommended to avoid using the .Rprofile file to set mc.cores (using the cores argument or setting mc.cores interactively or in a script is fine).
lpd_point	If calling stacking_weights() or pseudobma_weights() directly, a matrix of pointwise leave-one-out (or K-fold) log likelihoods evaluated for different models. It should be a N by K matrix where N is sample size and K is the number of models. Each column corresponds to one model. These values can be calculated approximately using loo() or by running exact leave-one-out or K-fold cross-validation.

Details

loo_model_weights() is a wrapper around the stacking_weights() and pseudobma_weights() functions that implements stacking, pseudo-BMA, and pseudo-BMA+ weighting for combining multiple predictive distributions. We can use approximate or exact leave-one-out cross-validation (LOO-CV) or K-fold CV to estimate the expected log predictive density (ELPD).

The stacking method (method="stacking"), which is the default for loo_model_weights(), combines all models by maximizing the leave-one-out predictive density of the combination distribution. That is, it finds the optimal linear combining weights for maximizing the leave-one-out log score.

The pseudo-BMA method (method="pseudobma") finds the relative weights proportional to the ELPD of each model. However, when method="pseudobma", the default is to also use the Bayesian bootstrap (BB=TRUE), which corresponds to the pseudo-BMA+ method. The Bayesian bootstrap takes into account the uncertainty of finite data points and regularizes the weights away from the extremes of 0 and 1.

In general, we recommend stacking for averaging predictive distributions, while pseudo-BMA+ can serve as a computationally easier alternative.

Value

A numeric vector containing one weight for each model.

References

- Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*. 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).
- Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)
- Yao, Y., Vehtari, A., Simpson, D., and Gelman, A. (2018) Using stacking to average Bayesian predictive distributions. *Bayesian Analysis*, advance publication, doi:10.1214/17-BA1091. ([online](#)).

See Also

- The **loo** package [vignettes](#), particularly [Bayesian Stacking and Pseudo-BMA weights using the loo package](#).
- [loo\(\)](#) for details on leave-one-out ELPD estimation.
- [constrOptim\(\)](#) for the choice of optimization methods and control-parameters.
- [relative_eff\(\)](#) for computing `r_eff`.

Examples

```
## Not run:
### Demonstrating usage after fitting models with RStan
library(rstan)

# generate fake data from N(0,1).
N <- 100
y <- rnorm(N, 0, 1)

# Suppose we have three models: N(-1, sigma), N(0.5, sigma) and N(0.6,sigma).
stan_code <- "
  data {
    int N;
    vector[N] y;
    real mu_fixed;
  }
  parameters {
    real<lower=0> sigma;
  }
  model {
    sigma ~ exponential(1);
    y ~ normal(mu_fixed, sigma);
  }
  generated quantities {
    vector[N] log_lik;
    for (n in 1:N) log_lik[n] = normal_lpdf(y[n]| mu_fixed, sigma);
  }"

mod <- stan_model(model_code = stan_code)
fit1 <- sampling(mod, data=list(N=N, y=y, mu_fixed=-1))
fit2 <- sampling(mod, data=list(N=N, y=y, mu_fixed=0.5))
```

```

fit3 <- sampling(mod, data=list(N=N, y=y, mu_fixed=0.6))
model_list <- list(fit1, fit2, fit3)
log_lik_list <- lapply(model_list, extract_log_lik)

# optional but recommended
r_eff_list <- lapply(model_list, function(x) {
  ll_array <- extract_log_lik(x, merge_chains = FALSE)
  relative_eff(exp(ll_array))
})

# stacking method:
wts1 <- loo_model_weights(
  log_lik_list,
  method = "stacking",
  r_eff_list = r_eff_list,
  optim_control = list(reltol=1e-10)
)
print(wts1)

# can also pass a list of psis_loo objects to avoid recomputing loo
loo_list <- lapply(1:length(log_lik_list), function(j) {
  loo(log_lik_list[[j]], r_eff = r_eff_list[[j]])
})

wts2 <- loo_model_weights(
  loo_list,
  method = "stacking",
  optim_control = list(reltol=1e-10)
)
all.equal(wts1, wts2)

# can provide names to be used in the results
loo_model_weights(setNames(loo_list, c("A", "B", "C")))

# pseudo-BMA+ method:
set.seed(1414)
loo_model_weights(loo_list, method = "pseudobma")

# pseudo-BMA method (set BB = FALSE):
loo_model_weights(loo_list, method = "pseudobma", BB = FALSE)

# calling stacking_weights or pseudobma_weights directly
lpd1 <- loo(log_lik_list[[1]], r_eff = r_eff_list[[1]])$pointwise[,1]
lpd2 <- loo(log_lik_list[[2]], r_eff = r_eff_list[[2]])$pointwise[,1]
lpd3 <- loo(log_lik_list[[3]], r_eff = r_eff_list[[3]])$pointwise[,1]
stacking_weights(cbind(lpd1, lpd2, lpd3))
pseudobma_weights(cbind(lpd1, lpd2, lpd3))
pseudobma_weights(cbind(lpd1, lpd2, lpd3), BB = FALSE)

## End(Not run)

```

loo_moment_match	<i>Moment matching for efficient approximate leave-one-out cross-validation (LOO)</i>
------------------	---

Description

Moment matching algorithm for updating a loo object when Pareto k estimates are large.

Usage

```
loo_moment_match(x, ...)

## Default S3 method:
loo_moment_match(
  x,
  loo,
  post_draws,
  log_lik_i,
  unconstrain_pars,
  log_prob_upars,
  log_lik_i_upars,
  max_iters = 30L,
  k_threshold = NULL,
  split = TRUE,
  cov = TRUE,
  cores = getOption("mc.cores", 1),
  ...
)
```

Arguments

x	A fitted model object.
...	Further arguments passed to the custom functions documented above.
loo	A loo object to be modified.
post_draws	A function the takes x as the first argument and returns a matrix of posterior draws of the model parameters.
log_lik_i	A function that takes x and i and returns a matrix (one column per chain) or a vector (all chains stacked) of log-likelihood draws of the i th observation based on the model x. If the draws are obtained using MCMC, the matrix with MCMC chains separated is preferred.
unconstrain_pars	A function that takes arguments x, and pars and returns posterior draws on the unconstrained space based on the posterior draws on the constrained space passed via pars.
log_prob_upars	A function that takes arguments x and upars and returns a matrix of log-posterior density values of the unconstrained posterior draws passed via upars.

log_lik_i_upars	A function that takes arguments <code>x</code> , <code>upars</code> , and <code>i</code> and returns a vector of log-likelihood draws of the i th observation based on the unconstrained posterior draws passed via <code>upars</code> .
max_iters	Maximum number of moment matching iterations. Usually this does not need to be modified. If the maximum number of iterations is reached, there will be a warning, and increasing <code>max_iters</code> may improve accuracy.
k_threshold	Threshold value for Pareto k values above which the moment matching algorithm is used. The default value is $\min(1 - 1/\log_{10}(S), 0.7)$, where S is the sample size.
split	Logical; Indicate whether to do the split transformation or not at the end of moment matching for each LOO fold.
cov	Logical; Indicate whether to match the covariance matrix of the samples or not. If FALSE, only the mean and marginal variances are matched.
cores	The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code> . The old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until <code>loo.cores</code> is removed in a future release. As of version 2.0.0 the default is now 1 core if <code>mc.cores</code> is not set , but we recommend using as many (or close to as many) cores as possible. • Note for Windows 10 users: it is strongly recommended to avoid using the <code>.Rprofile</code> file to set <code>mc.cores</code> (using the <code>cores</code> argument or setting <code>mc.cores</code> interactively or in a script is fine).

Details

The `loo_moment_match()` function is an S3 generic and we provide a default method that takes as arguments user-specified functions `post_draws`, `log_lik_i`, `unconstrain_pars`, `log_prob_upars`, and `log_lik_i_upars`. All of these functions should take `...` as an argument in addition to those specified for each function.

Value

The `loo_moment_match()` methods return an updated loo object. The structure of the updated loo object is similar, but the method also stores the original Pareto k diagnostic values in the `diagnostics` field.

Methods (by class)

- `loo_moment_match(default)`: A default method that takes as arguments a user-specified model object `x`, a loo object and user-specified functions `post_draws`, `log_lik_i`, `unconstrain_pars`, `log_prob_upars`, and `log_lik_i_upars`.

References

Paananen, T., Piironen, J., Buerkner, P.-C., Vehtari, A. (2021). Implicitly adaptive importance sampling. *Statistics and Computing*, 31, 16. doi:10.1007/s11222-020-09982-2. arXiv preprint arXiv:1906.08850.

See Also

[loo\(\)](#), [loo_moment_match_split\(\)](#)

Examples

```
# See the vignette for loo_moment_match()
```

```
loo_moment_match_split
```

Split moment matching for efficient approximate leave-one-out cross-validation (LOO)

Description

A function that computes the split moment matching importance sampling loo. Takes in the moment matching total transformation, transforms only half of the draws, and computes a single elpd using multiple importance sampling.

Usage

```
loo_moment_match_split(
  x,
  upars,
  cov,
  total_shift,
  total_scaling,
  total_mapping,
  i,
  log_prob_upars,
  log_lik_i_upars,
  r_eff_i,
  cores,
  is_method,
  ...
)
```

Arguments

<code>x</code>	A fitted model object.
<code>upars</code>	A matrix containing the model parameters in unconstrained space where they can have any real value.
<code>cov</code>	Logical; Indicate whether to match the covariance matrix of the samples or not. If FALSE, only the mean and marginal variances are matched.
<code>total_shift</code>	A vector representing the total shift made by the moment matching algorithm.
<code>total_scaling</code>	A vector representing the total scaling of marginal variance made by the moment matching algorithm.

total_mapping	A vector representing the total covariance transformation made by the moment matching algorithm.
i	Observation index.
log_prob_upars	A function that takes arguments <code>x</code> and <code>upars</code> and returns a matrix of log-posterior density values of the unconstrained posterior draws passed via <code>upars</code> .
log_lik_i_upars	A function that takes arguments <code>x</code> , <code>upars</code> , and <code>i</code> and returns a vector of log-likelihood draws of the <code>i</code> th observation based on the unconstrained posterior draws passed via <code>upars</code> .
r_eff_i	MCMC relative effective sample size of the <code>i</code> 'th log likelihood draws.
cores	The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code> . The old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until <code>loo.cores</code> is removed in a future release. As of version 2.0.0 the default is now 1 core if <code>mc.cores</code> is not set , but we recommend using as many (or close to as many) cores as possible. • Note for Windows 10 users: it is strongly recommended to avoid using the <code>.Rprofile</code> file to set <code>mc.cores</code> (using the <code>cores</code> argument or setting <code>mc.cores</code> interactively or in a script is fine).
is_method	The importance sampling method to use. The following methods are implemented: • <code>"psis"</code>: Pareto-Smoothed Importance Sampling (PSIS). Default method. • <code>"tis"</code>: Truncated Importance Sampling (TIS) with truncation at <code>sqrt(S)</code>, where <code>S</code> is the number of posterior draws. • <code>"sis"</code>: Standard Importance Sampling (SIS).
...	Further arguments passed to the custom functions documented above.

Value

A list containing the updated log-importance weights and log-likelihood values. Also returns the updated MCMC effective sample size and the integrand-specific log-importance weights.

References

Paananen, T., Piironen, J., Buerkner, P.-C., Vehtari, A. (2021). Implicitly adaptive importance sampling. *Statistics and Computing*, 31, 16. doi:10.1007/s11222-020-09982-2. arXiv preprint arXiv:1906.08850.

See Also

[loo\(\)](#), [loo_moment_match\(\)](#)

`loo_predictive_metric` *Estimate leave-one-out predictive performance..*

Description

The `loo_predictive_metric()` function computes estimates of leave-one-out predictive metrics given a set of predictions and observations. Currently supported metrics are mean absolute error, mean squared error and root mean squared error for continuous predictions and accuracy and balanced accuracy for binary classification. Predictions are passed on to the [E_loo\(\)](#) function, so this function assumes that the PSIS approximation is working well.

Usage

```
loo_predictive_metric(x, ...)

## S3 method for class 'matrix'
loo_predictive_metric(
  x,
  y,
  log_lik,
  ...,
  metric = c("mae", "rmse", "mse", "acc", "balanced_acc"),
  r_eff = 1,
  cores = getOption("mc.cores", 1)
)
```

Arguments

<code>x</code>	A numeric matrix of predictions.
<code>...</code>	Additional arguments passed on to E_loo()
<code>y</code>	A numeric vector of observations. Length should be equal to the number of rows in <code>x</code> .
<code>log_lik</code>	A matrix of pointwise log-likelihoods. Should be of same dimension as <code>x</code> .
<code>metric</code>	The type of predictive metric to be used. Currently supported options are "mae", "rmse" and "mse" for regression and for binary classification "acc" and "balanced_acc". "mae" Mean absolute error. "mse" Mean squared error. "rmse" Root mean squared error, given by as the square root of MSE. "acc" The proportion of predictions indicating the correct outcome. "balanced_acc" Balanced accuracy is given by the average of true positive and true negative rates.
<code>r_eff</code>	A Vector of relative effective sample size estimates containing one element per observation. See psis() for more details.
<code>cores</code>	The number of cores to use for parallelization of <code>[psis()]</code> . See psis() for details.

Value

A list with the following components:

`estimate` Estimate of the given metric.

`se` Standard error of the estimate.

Examples

```
if (requireNamespace("rstanarm", quietly = TRUE)) {
  # Use rstanarm package to quickly fit a model and get both a log-likelihood
  # matrix and draws from the posterior predictive distribution
  library("rstanarm")

  # data from help("lm")
  ctl <- c(4.17, 5.58, 5.18, 6.11, 4.50, 4.61, 5.17, 4.53, 5.33, 5.14)
  trt <- c(4.81, 4.17, 4.41, 3.59, 5.87, 3.83, 6.03, 4.89, 4.32, 4.69)
  d <- data.frame(
    weight = c(ctl, trt),
    group = gl(2, 10, 20, labels = c("Ctl", "Trt"))
  )
  fit <- stan_glm(weight ~ group, data = d, refresh = 0)
  ll <- log_lik(fit)
  r_eff <- relative_eff(exp(-ll), chain_id = rep(1:4, each = 1000))

  mu_pred <- posterior_epred(fit)
  # Leave-one-out mean absolute error of predictions
  mae <- loo_predictive_metric(x = mu_pred, y = d$weight, log_lik = ll,
    pred_error = 'mae', r_eff = r_eff)
  # Leave-one-out 90%-quantile of mean absolute error
  mae_90q <- loo_predictive_metric(x = mu_pred, y = d$weight, log_lik = ll,
    pred_error = 'mae', r_eff = r_eff,
    type = 'quantile', probs = 0.9)
}
```

loo_subsample

Efficient approximate LOO-CV using subsampling

Description

Efficiently approximates leave-one-out cross-validation (LOO-CV) via subsampling. Less costly and more approximate computation is made for all LOO-fold, and more costly and accurate computations are made only for $m < N$ LOO-folds.

Usage

```
loo_subsample(x, ...)
```

```
## S3 method for class ``function``
```

```

loo_subsample(
  x,
  ...,
  data = NULL,
  draws = NULL,
  observations = 400,
  log_p = NULL,
  log_g = NULL,
  r_eff = 1,
  save_psis = FALSE,
  cores = getOption("mc.cores", 1),
  loo_approximation = "plpd",
  loo_approximation_draws = NULL,
  estimator = "diff_srs",
  llgrad = NULL,
  llhess = NULL
)

```

Arguments

<code>x</code>	A function. The Methods (by class) section, below, has detailed descriptions of how to specify the inputs.
<code>data, draws, ...</code>	For <code>loo_subsample.function()</code> , these are the data, posterior draws, and other arguments to pass to the log-likelihood function. Note that for some <code>loo_approximations</code> , the draws will be replaced by the posteriors summary statistics to compute loo approximations. See argument <code>loo_approximation</code> for details.
<code>observations</code>	The subsample observations to use. The argument can take four (4) types of arguments: • <code>NULL</code> to use all observations. The algorithm then just uses standard <code>loo()</code> or <code>loo_approximate_posterior()</code>. • A single integer to specify the number of observations to be subsampled. • A vector of integers to provide the indices used to subset the data. <i>These observations need to be subsampled with the same scheme as given by the estimator argument.</i> • A <code>psis_loo_ss</code> object to use the same observations that were used in a previous call to <code>loo_subsample()</code>.
<code>log_p, log_g</code>	Should be supplied only if approximate posterior draws are used. The default (<code>NULL</code>) indicates draws are from "true" posterior (i.e. using MCMC). If not <code>NULL</code> then they should be specified as described in loo_approximate_posterior() .
<code>r_eff</code>	Vector of relative effective sample size estimates for the likelihood ($\exp(\log_{\text{lik}})$) of each observation. This is related to the relative efficiency of estimating the normalizing term in self-normalized importance sampling when using posterior draws obtained with MCMC. If MCMC draws are used and <code>r_eff</code> is not provided then the reported PSIS effective sample sizes and Monte Carlo error estimates can be over-optimistic. If the posterior draws are (near) independent then <code>r_eff=1</code> can be used. <code>r_eff</code> has to be a scalar (same value is used for all observations) or a vector with length equal to the number of observations. The

	default value is 1. See the <code>relative_eff()</code> helper functions for help computing <code>r_eff</code> .
<code>save_psis</code>	Should the "psis" object created internally by <code>loo_subsample()</code> be saved in the returned object? See <code>loo()</code> for details.
<code>cores</code>	The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code>. The old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until <code>loo.cores</code> is removed in a future release. As of version 2.0.0 the default is now 1 core if <code>mc.cores</code> is not set, but we recommend using as many (or close to as many) cores as possible. • Note for Windows 10 users: it is strongly recommended to avoid using the <code>.Rprofile</code> file to set <code>mc.cores</code> (using the <code>cores</code> argument or setting <code>mc.cores</code> interactively or in a script is fine).
<code>loo_approximation</code>	What type of approximation of the <code>loo_i</code>'s should be used? The default is "plpd" (the log predictive density using the posterior expectation). There are six different methods implemented to approximate <code>loo_i</code>'s (see the references for more details): • "plpd": uses the lpd based on point estimates (i.e., $p(y_i \hat{\theta})$). • "lpd": uses the lpds (i.e., $p(y_i y)$). • "tis": uses truncated importance sampling to approximate PSIS-LOO. • "waic": uses waic (i.e., $p(y_i y) - p_{waic}$). • "waic_grad_marginal": uses waic approximation using first order delta method and posterior marginal variances to approximate p_{waic} (ie. $p(y_i \hat{\theta}) - p_{waic_grad_marginal}$). Requires gradient of likelihood function. • "waic_grad": uses waic approximation using first order delta method and posterior covariance to approximate p_{waic} (ie. $p(y_i \hat{\theta}) - p_{waic_grad}$). Requires gradient of likelihood function. • "waic_hess": uses waic approximation using second order delta method and posterior covariance to approximate p_{waic} (ie. $p(y_i \hat{\theta}) - p_{waic_grad}$). Requires gradient and Hessian of likelihood function. As point estimates of $\hat{\theta}$, the posterior expectations of the parameters are used.
<code>loo_approximation_draws</code>	The number of posterior draws used when integrating over the posterior. This is used if <code>loo_approximation</code> is set to "lpd", "waic", or "tis".
<code>estimator</code>	How should <code>elpd_loo</code>, <code>p_loo</code> and <code>looic</code> be estimated? The default is "diff_srs". • "diff_srs": uses the difference estimator with simple random sampling without replacement (srs). <code>p_loo</code> is estimated using standard srs. (Magnusson et al., 2020) • "hh": uses the Hansen-Hurwitz estimator with sampling with replacement proportional to size, where <code>abs</code> of <code>loo_approximation</code> is used as size. (Magnusson et al., 2019) • "srs": uses simple random sampling and ordinary estimation.
<code>llgrad</code>	The gradient of the log-likelihood. This is only used when <code>loo_approximation</code> is "waic_grad", "waic_grad_marginal", or "waic_hess". The default is NULL.

`llhess` The Hessian of the log-likelihood. This is only used with `loo_approximation = "waic_hess"`. The default is `NULL`.

Details

The `loo_subsample()` function is an S3 generic and a methods is currently provided for log-likelihood functions. The implementation works for both MCMC and for posterior approximations where it is possible to compute the log density for the approximation.

Value

`loo_subsample()` returns a named list with class `c("psis_loo_ss", "psis_loo", "loo")`. This has the same structure as objects returned by `loo()` but with the additional slot:

- `loo_subsampling`: A list with two vectors, `log_p` and `log_g`, of the same length containing the posterior density and the approximation density for the individual draws.

Methods (by class)

- `loo_subsample(`function`)`: A function `f()` that takes arguments `data_i` and `draws` and returns a vector containing the log-likelihood for a single observation `i` evaluated at each posterior draw. The function should be written such that, for each observation `i` in `1:N`, evaluating

```
f(data_i = data[i,, drop=FALSE], draws = draws)
```

results in a vector of length `S` (size of posterior sample). The log-likelihood function can also have additional arguments but `data_i` and `draws` are required.

If using the function method then the arguments `data` and `draws` must also be specified in the call to `loo()`:

- `data`: A data frame or matrix containing the data (e.g. observed outcome and predictors) needed to compute the pointwise log-likelihood. For each observation `i`, the `i`th row of `data` will be passed to the `data_i` argument of the log-likelihood function.
- `draws`: An object containing the posterior draws for any parameters needed to compute the pointwise log-likelihood. Unlike `data`, which is indexed by observation, for each observation the entire object `draws` will be passed to the `draws` argument of the log-likelihood function.
- The `...` can be used if your log-likelihood function takes additional arguments. These arguments are used like the `draws` argument in that they are recycled for each observation.

References

Magnusson, M., Riis Andersen, M., Jonasson, J. and Vehtari, A. (2019). Leave-One-Out Cross-Validation for Large Data. In *Thirty-sixth International Conference on Machine Learning*, PMLR 97:4244-4253.

Magnusson, M., Riis Andersen, M., Jonasson, J. and Vehtari, A. (2020). Leave-One-Out Cross-Validation for Model Comparison in Large Data. In *Proceedings of the 23rd International Conference on Artificial Intelligence and Statistics (AISTATS)*, PMLR 108:341-351.

See Also

[loo\(\)](#), [psis\(\)](#), [loo_compare\(\)](#)

nobs.psis_loo_ss	<i>The number of observations in a psis_loo_ss object.</i>
------------------	--

Description

The number of observations in a psis_loo_ss object.

Usage

```
## S3 method for class 'psis_loo_ss'
nobs(object, ...)
```

Arguments

object	a psis_loo_ss object.
...	Currently unused.

obs_idx	<i>Get observation indices used in subsampling</i>
---------	--

Description

Get observation indices used in subsampling

Usage

```
obs_idx(x, rep = TRUE)
```

Arguments

x	A psis_loo_ss object.
rep	If sampling with replacement is used, an observation can have multiple samples and these are then repeated in the returned object if rep=TRUE (e.g., a vector c(1, 1, 2) indicates that observation 1 has been subsampled two times). If rep=FALSE only the unique indices are returned.

Value

An integer vector.

pareto-k-diagnostic *Diagnostics for Pareto smoothed importance sampling (PSIS)*

Description

Print a diagnostic table summarizing the estimated Pareto shape parameters and PSIS effective sample sizes, find the indexes of observations for which the estimated Pareto shape parameter k is larger than some threshold value, or plot observation indexes vs. diagnostic estimates. The **Details** section below provides a brief overview of the diagnostics, but we recommend consulting Vehtari, Gelman, and Gabry (2017) and Vehtari, Simpson, Gelman, Yao, and Gabry (2024) for full details.

Usage

```
pareto_k_table(x)

pareto_k_ids(x, threshold = NULL)

pareto_k_values(x)

pareto_k_influence_values(x)

psis_n_eff_values(x)

mcse_loo(x, threshold = NULL)

## S3 method for class 'psis_loo'
plot(
  x,
  diagnostic = c("k", "ESS", "n_eff"),
  ...,
  label_points = FALSE,
  main = "PSIS diagnostic plot"
)

## S3 method for class 'psis'
plot(
  x,
  diagnostic = c("k", "ESS", "n_eff"),
  ...,
  label_points = FALSE,
  main = "PSIS diagnostic plot"
)
```

Arguments

`x` An object created by `loo()` or `psis()`.

threshold	For <code>pareto_k_ids()</code> , threshold is the minimum k value to flag (default is a sample size S dependent threshold $1 - 1 / \log_{10}(S)$). For <code>mcse_loo()</code> , if any k estimates are greater than threshold the MCSE estimate is returned as NA See Details for the motivation behind these defaults.
diagnostic	For the <code>plot</code> method, which diagnostic should be plotted? The options are "k" for Pareto k estimates (the default), or "ESS" or "n_eff" for PSIS effective sample size estimates.
label_points, ...	For the <code>plot()</code> method, if <code>label_points</code> is TRUE the observation numbers corresponding to any values of k greater than the diagnostic threshold will be displayed in the plot. Any arguments specified in ... will be passed to <code>graphics::text()</code> and can be used to control the appearance of the labels.
main	For the <code>plot()</code> method, a title for the plot.

Details

The reliability and approximate convergence rate of the PSIS-based estimates can be assessed using the estimates for the shape parameter k of the generalized Pareto distribution. The diagnostic threshold for Pareto k depends on sample size S (sample size dependent threshold was introduced by Vehtari et al. (2024), and before that fixed thresholds of 0.5 and 0.7 were recommended). For simplicity, loo package uses the nominal sample size S when computing the sample size specific threshold. This provides an optimistic threshold if the effective sample size is less than 2200, but if MCMC-ESS > $S/2$ the difference is usually negligible. Thinning of MCMC draws can be used to improve the ratio ESS/ S .

- If $k < \min(1 - 1/\log_{10}(S), 0.7)$, where S is the sample size, the PSIS estimate and the corresponding Monte Carlo standard error estimate are reliable.
- If $1 - 1/\log_{10}(S) \leq k < 0.7$, the PSIS estimate and the corresponding Monte Carlo standard error estimate are not reliable, but increasing the (effective) sample size S above 2200 may help (this will increase the sample size specific threshold $(1 - 1/\log_{10}(2200)) > 0.7$ and then the bias specific threshold 0.7 dominates).
- If $0.7 \leq k < 1$, the PSIS estimate and the corresponding Monte Carlo standard error have large bias and are not reliable. Increasing the sample size may reduce the variability in k estimate, which may result in lower k estimate, too.
- If $k \geq 1$, the target distribution is estimated to have a non-finite mean. The PSIS estimate and the corresponding Monte Carlo standard error are not well defined. Increasing the sample size may reduce the variability in the k estimate, which may also result in a lower k estimate.

What if the estimated tail shape parameter k exceeds the diagnostic threshold?: Importance sampling is likely to work less well if the marginal posterior $p(\theta^s|y)$ and LOO posterior $p(\theta^s|y_{-i})$ are very different, which is more likely to happen with a non-robust model and highly influential observations. If the estimated tail shape parameter k exceeds the diagnostic threshold, the user should be warned. (Note: If k is greater than the diagnostic threshold then WAIC is also likely to fail, but WAIC lacks as accurate diagnostic.) When using PSIS in the context of approximate LOO-CV, we recommend one of the following actions:

- With some additional computations, it is possible to transform the MCMC draws from the posterior distribution to obtain more reliable importance sampling estimates. This results in a

smaller shape parameter k . See `loo_moment_match()` and the vignette *Avoiding model refits in leave-one-out cross-validation with moment matching* for an example of this.

- Sampling from a leave-one-out mixture distribution (see the vignette *Mixture IS leave-one-out cross-validation for high-dimensional Bayesian models*), directly from $p(\theta^s | y_{-i})$ for the problematic observations i , or using K -fold cross-validation (see the vignette *Holdout validation and K-fold cross-validation of Stan programs with the loo package*) will generally be more stable.
- Using a model that is more robust to anomalous observations will generally make approximate LOO-CV more stable.

Observation influence statistics: The estimated shape parameter k for each observation can be used as a measure of the observation's influence on posterior distribution of the model. These can be obtained with `pareto_k_influence_values()`.

Effective sample size and error estimates: In the case that we obtain the samples from the proposal distribution via MCMC the `loo` package also computes estimates for the Monte Carlo error and the effective sample size for importance sampling, which are more accurate for PSIS than for IS and TIS (see Vehtari et al (2024) for details). However, the PSIS effective sample size estimate will be **over-optimistic when the estimate of k is greater than $\min(1 - 1/\log_{10}(S), 0.7)$** , where S is the sample size.

Value

`pareto_k_table()` returns an object of class "pareto_k_table", which is a matrix with columns "Count", "Proportion", and "Min. n_eff", and has its own print method.

`pareto_k_ids()` returns an integer vector indicating which observations have Pareto k estimates above threshold.

`pareto_k_values()` returns a vector of the estimated Pareto k parameters. These represent the reliability of sampling.

`pareto_k_influence_values()` returns a vector of the estimated Pareto k parameters. These represent influence of the observations on the model posterior distribution.

`psis_n_eff_values()` returns a vector of the estimated PSIS effective sample sizes.

`mcse_loo()` returns the Monte Carlo standard error (MCSE) estimate for PSIS-LOO. MCSE will be NA if any Pareto k values are above threshold.

The `plot()` method is called for its side effect and does not return anything. If `x` is the result of a call to `loo()` or `psis()` then `plot(x, diagnostic)` produces a plot of the estimates of the Pareto shape parameters (`diagnostic = "k"`) or estimates of the PSIS effective sample sizes (`diagnostic = "ESS"`).

References

- Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*. 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).
- Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)

See Also

- [psis\(\)](#) for the implementation of the PSIS algorithm.
- The [FAQ page](#) on the **loo** website for answers to frequently asked questions.

pointwise

*Convenience function for extracting pointwise estimates***Description**

Convenience function for extracting pointwise estimates

Usage

```
pointwise(x, estimate, ...)
```

```
## S3 method for class 'loo'
pointwise(x, estimate, ...)
```

Arguments

<code>x</code>	A loo object, for example one returned by loo() , loo_subsample() , loo_approximate_posterior() , loo_moment_match() , etc.
<code>estimate</code>	Which pointwise estimate to return. By default all are returned. The objects returned by the different functions (loo() , loo_subsample() , etc.) have slightly different estimates available. Typically at a minimum the estimates <code>elpd_loo</code> , <code>looic</code> , <code>mcse_elpd_loo</code> , <code>p_loo</code> , and <code>influence_pareto_k</code> will be available, but there may be others.
<code>...</code>	Currently ignored.

Value

A vector of length equal to the number of observations.

Examples

```
x <- loo(example_loglik_array())
pointwise(x, "elpd_loo")
```

print.loo

Print methods

Description

Print methods

Usage

```
## S3 method for class 'loo'
print(x, digits = 1, ...)

## S3 method for class 'waic'
print(x, digits = 1, ...)

## S3 method for class 'psis_loo'
print(x, digits = 1, plot_k = FALSE, ...)

## S3 method for class 'importance_sampling_loo'
print(x, digits = 1, plot_k = FALSE, ...)

## S3 method for class 'psis_loo_ap'
print(x, digits = 1, plot_k = FALSE, ...)

## S3 method for class 'psis'
print(x, digits = 1, plot_k = FALSE, ...)

## S3 method for class 'importance_sampling'
print(x, digits = 1, plot_k = FALSE, ...)

## S3 method for class 'kfold'
print(x, digits = 1, plot_k = FALSE, ...)
```

Arguments

x	An object returned by <code>loo()</code> , <code>psis()</code> , or <code>waic()</code> .
digits	An integer passed to <code>base::round()</code> .
...	Arguments passed to <code>plot.psis_loo()</code> if <code>plot_k</code> is TRUE.
plot_k	Logical. If TRUE the estimates of the Pareto shape parameter k are plotted. Ignored if <code>x</code> was generated by <code>waic()</code> . To just plot k without printing use the <code>plot()</code> method for 'loo' objects.

Value

`x`, invisibly.

See Also[pareto-k-diagnostic](#)

psis*Pareto smoothed importance sampling (PSIS)*

Description

Implementation of Pareto smoothed importance sampling (PSIS), a method for stabilizing importance ratios. The version of PSIS implemented here corresponds to the algorithm presented in Vehtari, Simpson, Gelman, Yao, and Gabry (2024). For PSIS diagnostics see the [pareto-k-diagnostic](#) page.

Usage

```
psis(log_ratios, ...)

## S3 method for class 'array'
psis(log_ratios, ..., r_eff = 1, cores = getOption("mc.cores", 1))

## S3 method for class 'matrix'
psis(log_ratios, ..., r_eff = 1, cores = getOption("mc.cores", 1))

## Default S3 method:
psis(log_ratios, ..., r_eff = 1)

is.psis(x)

is.sis(x)

is.tis(x)
```

Arguments

<code>log_ratios</code>	An array, matrix, or vector of importance ratios on the log scale (for PSIS-LOO these are <i>negative</i> log-likelihood values). See the Methods (by class) section below for a detailed description of how to specify the inputs for each method.
<code>...</code>	Arguments passed on to the various methods.
<code>r_eff</code>	Vector of relative effective sample size estimates containing one element per observation. The values provided should be the relative effective sample sizes of $1/\exp(\log_ratios)$ (i.e., $1/ratios$). This is related to the relative efficiency of estimating the normalizing term in self-normalizing importance sampling. If <code>r_eff</code> is not provided then the reported PSIS effective sample sizes and Monte Carlo error estimates can be over-optimistic. If the posterior draws are (near) independent then <code>r_eff=1</code> can be used. <code>r_eff</code> has to be a scalar (same value

is used for all observations) or a vector with length equal to the number of observations. The default value is 1. See the [relative_eff\(\)](#) helper function for computing `r_eff`.

cores	The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code>. The old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until <code>loo.cores</code> is removed in a future release. As of version 2.0.0 the default is now 1 core if <code>mc.cores</code> is not set, but we recommend using as many (or close to as many) cores as possible. • Note for Windows 10 users: it is strongly recommended to avoid using the <code>.Rprofile</code> file to set <code>mc.cores</code> (using the <code>cores</code> argument or setting <code>mc.cores</code> interactively or in a script is fine).
x	For <code>is.psis()</code> , an object to check.

Value

The `psis()` methods return an object of class "psis", which is a named list with the following components:

log_weights Vector or matrix of smoothed (and truncated) but *unnormalized* log weights. To get normalized weights use the [weights\(\)](#) method provided for objects of class "psis".

diagnostics A named list containing two vectors:

- **pareto_k**: Estimates of the shape parameter k of the generalized Pareto distribution. See the [pareto-k-diagnostic](#) page for details.
- **n_eff**: PSIS effective sample size estimates.

Objects of class "psis" also have the following [attributes](#):

norm_const_log Vector of precomputed values of `colLogSumExps(log_weights)` that are used internally by the `weights` method to normalize the log weights.

tail_len Vector of tail lengths used for fitting the generalized Pareto distribution.

r_eff If specified, the user's `r_eff` argument.

dims Integer vector of length 2 containing S (posterior sample size) and N (number of observations).

method Method used for importance sampling, here `psis`.

Methods (by class)

- **psis(array)**: An I by C by N array, where I is the number of MCMC iterations per chain, C is the number of chains, and N is the number of data points.
- **psis(matrix)**: An S by N matrix, where S is the size of the posterior sample (with all chains merged) and N is the number of data points.
- **psis(default)**: A vector of length S (posterior sample size).

References

Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).

Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)

See Also

- [loo\(\)](#) for approximate LOO-CV using PSIS.
- [pareto-k-diagnostic](#) for PSIS diagnostics.
- The **loo** package [vignettes](#) for demonstrations.
- The [FAQ page](#) on the **loo** website for answers to frequently asked questions.

Examples

```
log_ratios <- -1 * example_loglik_array()
r_eff <- relative_eff(exp(-log_ratios))
psis_result <- psis(log_ratios, r_eff = r_eff)
str(psis_result)
plot(psis_result)

# extract smoothed weights
lw <- weights(psis_result) # default args are log=TRUE, normalize=TRUE
ulw <- weights(psis_result, normalize=FALSE) # unnormalized log-weights

w <- weights(psis_result, log=FALSE) # normalized weights (not log-weights)
uw <- weights(psis_result, log=FALSE, normalize = FALSE) # unnormalized weights
```

psislw

Pareto smoothed importance sampling (deprecated, old version)

Description

As of version 2.0.0 this function is **deprecated**. Please use the [psis\(\)](#) function for the new PSIS algorithm.

Usage

```
psislw(
  lw,
  wcp = 0.2,
  wtrunc = 3/4,
  cores = getOption("mc.cores", 1),
```

```

    llfun = NULL,
    llargs = NULL,
    ...
)

```

Arguments

lw	A matrix or vector of log weights. For computing LOO, $lw = -\log_lik$, the <i>negative</i> of an S (simulations) by N (data points) pointwise log-likelihood matrix.
wcp	The proportion of importance weights to use for the generalized Pareto fit. The $100 \times wcp\%$ from which to estimate the parameters of the generalized Pareto distribution.
wtrunc	For truncating very large weights to S^{wtrunc} . Set to zero for no truncation.
cores	The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code> , the old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until it is removed. As of version 2.0.0, the default is now 1 core if <code>mc.cores</code> is not set, but we recommend using as many (or close to as many) cores as possible.
llfun, llargs	See <code>loo.function()</code> .
...	Ignored when <code>psislw()</code> is called directly. The ... is only used internally when <code>psislw()</code> is called by the <code>loo()</code> function.

Value

A named list with components `lw_smooth` (modified log weights) and `pareto_k` (estimated generalized Pareto shape parameter(s) k).

References

- Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).
- Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)

See Also

[pareto-k-diagnostic](#) for PSIS diagnostics.

relative_eff	<i>Convenience function for computing relative efficiencies</i>
--------------	---

Description

relative_eff() computes the the MCMC effective sample size divided by the total sample size.

Usage

```
relative_eff(x, ...)

## Default S3 method:
relative_eff(x, chain_id, ...)

## S3 method for class 'matrix'
relative_eff(x, chain_id, ..., cores = getOption("mc.cores", 1))

## S3 method for class 'array'
relative_eff(x, ..., cores = getOption("mc.cores", 1))

## S3 method for class '`function`'
relative_eff(
  x,
  chain_id,
  ...,
  cores = getOption("mc.cores", 1),
  data = NULL,
  draws = NULL
)

## S3 method for class 'importance_sampling'
relative_eff(x, ...)
```

Arguments

x	A vector, matrix, 3-D array, or function. See the Methods (by class) section below for details on specifying x, but where "log-likelihood" is mentioned replace it with one of the following depending on the use case: • For use with the <code>loo()</code> function, the values in x (or generated by x, if a function) should be likelihood values (i.e., $\exp(\log_lik)$), not on the log scale. • For generic use with <code>psis()</code>, the values in x should be the reciprocal of the importance ratios (i.e., $\exp(-\log_ratios)$).
chain_id	A vector of length <code>NROW(x)</code> containing MCMC chain indexes for each each row of x (if a matrix) or each value in x (if a vector). No chain_id is needed if x is a 3-D array. If there are C chains then valid chain indexes are values in 1:C.

cores The number of cores to use for parallelization.
 data, draws, ... Same as for the `loo()` function method.

Value

A vector of relative effective sample sizes.

Methods (by class)

- `relative_eff(default)`: A vector of length S (posterior sample size).
- `relative_eff(matrix)`: An S by N matrix, where S is the size of the posterior sample (with all chains merged) and N is the number of data points.
- `relative_eff(array)`: An I by C by N array, where I is the number of MCMC iterations per chain, C is the number of chains, and N is the number of data points.
- `relative_eff(`function`)`: A function $f()$ that takes arguments `data_i` and `draws` and returns a vector containing the log-likelihood for a single observation i evaluated at each posterior draw. The function should be written such that, for each observation i in $1:N$, evaluating

```
f(data_i = data[i,, drop=FALSE], draws = draws)
```

results in a vector of length S (size of posterior sample). The log-likelihood function can also have additional arguments but `data_i` and `draws` are required.

If using the function method then the arguments `data` and `draws` must also be specified in the call to `loo()`:

- `data`: A data frame or matrix containing the data (e.g. observed outcome and predictors) needed to compute the pointwise log-likelihood. For each observation i , the i th row of `data` will be passed to the `data_i` argument of the log-likelihood function.
- `draws`: An object containing the posterior draws for any parameters needed to compute the pointwise log-likelihood. Unlike `data`, which is indexed by observation, for each observation the entire object `draws` will be passed to the `draws` argument of the log-likelihood function.
- The ... can be used if your log-likelihood function takes additional arguments. These arguments are used like the `draws` argument in that they are recycled for each observation.
- `relative_eff(importance_sampling)`: If x is an object of class "psis", `relative_eff()` simply returns the `r_eff` attribute of x .

Examples

```
LLarr <- example_loglik_array()
LLmat <- example_loglik_matrix()
dim(LLarr)
dim(LLmat)

rel_n_eff_1 <- relative_eff(exp(LLarr))
rel_n_eff_2 <- relative_eff(exp(LLmat), chain_id = rep(1:2, each = 500))
all.equal(rel_n_eff_1, rel_n_eff_2)
```

sis

Standard importance sampling (SIS)

Description

Implementation of standard importance sampling (SIS).

Usage

```
sis(log_ratios, ...)

## S3 method for class 'array'
sis(log_ratios, ..., r_eff = NULL, cores = getOption("mc.cores", 1))

## S3 method for class 'matrix'
sis(log_ratios, ..., r_eff = NULL, cores = getOption("mc.cores", 1))

## Default S3 method:
sis(log_ratios, ..., r_eff = NULL)
```

Arguments

- | | |
|------------|---|
| log_ratios | An array, matrix, or vector of importance ratios on the log scale (for Importance sampling LOO, these are <i>negative</i> log-likelihood values). See the Methods (by class) section below for a detailed description of how to specify the inputs for each method. |
| ... | Arguments passed on to the various methods. |
| r_eff | Vector of relative effective sample size estimates containing one element per observation. The values provided should be the relative effective sample sizes of $1/\exp(\log_ratios)$ (i.e., $1/ratios$). This is related to the relative efficiency of estimating the normalizing term in self-normalizing importance sampling. See the relative_eff() helper function for computing <code>r_eff</code> . If using <code>psis</code> with draws of the <code>log_ratios</code> not obtained from MCMC then the warning message thrown when not specifying <code>r_eff</code> can be disabled by setting <code>r_eff</code> to NA. |
| cores | The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code> . The old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until <code>loo.cores</code> is removed in a future release. As of version 2.0.0 the default is now 1 core if <code>mc.cores</code> is not set , but we recommend using as many (or close to as many) cores as possible. <ul style="list-style-type: none"> • Note for Windows 10 users: it is strongly recommended to avoid using the <code>.Rprofile</code> file to set <code>mc.cores</code> (using the <code>cores</code> argument or setting <code>mc.cores</code> interactively or in a script is fine). |

Value

The `sis()` methods return an object of class "sis", which is a named list with the following components:

`log_weights` Vector or matrix of smoothed but *unnormalized* log weights. To get normalized weights use the `weights()` method provided for objects of class `sis`.

`diagnostics` A named list containing one vector:

- `pareto_k`: Not used in `sis`, all set to 0.
- `n_eff`: effective sample size estimates.

Objects of class "sis" also have the following [attributes](#):

`norm_const_log` Vector of precomputed values of `colLogSumExps(log_weights)` that are used internally by the `weights` method to normalize the log weights.

`r_eff` If specified, the user's `r_eff` argument.

`tail_len` Not used for `sis`.

`dims` Integer vector of length 2 containing S (posterior sample size) and N (number of observations).

`method` Method used for importance sampling, here `sis`.

Methods (by class)

- `sis(array)`: An I by C by N array, where I is the number of MCMC iterations per chain, C is the number of chains, and N is the number of data points.
- `sis(matrix)`: An S by N matrix, where S is the size of the posterior sample (with all chains merged) and N is the number of data points.
- `sis(default)`: A vector of length S (posterior sample size).

References

Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*, 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).

Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)

See Also

- [psis\(\)](#) for approximate LOO-CV using PSIS.
- [loo\(\)](#) for approximate LOO-CV.
- [pareto-k-diagnostic](#) for PSIS diagnostics.

Examples

```
log_ratios <- -1 * example_loglik_array()
r_eff <- relative_eff(exp(-log_ratios))
sis_result <- sis(log_ratios, r_eff = r_eff)
str(sis_result)

# extract smoothed weights
lw <- weights(sis_result) # default args are log=TRUE, normalize=TRUE
ulw <- weights(sis_result, normalize=FALSE) # unnormalized log-weights

w <- weights(sis_result, log=FALSE) # normalized weights (not log-weights)
uw <- weights(sis_result, log=FALSE, normalize = FALSE) # unnormalized weights
```

tis

Truncated importance sampling (TIS)

Description

Implementation of truncated (self-normalized) importance sampling (TIS), truncated at $S^{(1/2)}$ as recommended by Ionides (2008).

Usage

```
tis(log_ratios, ...)

## S3 method for class 'array'
tis(log_ratios, ..., r_eff = 1, cores = getOption("mc.cores", 1))

## S3 method for class 'matrix'
tis(log_ratios, ..., r_eff = 1, cores = getOption("mc.cores", 1))

## Default S3 method:
tis(log_ratios, ..., r_eff = 1)
```

Arguments

log_ratios	An array, matrix, or vector of importance ratios on the log scale (for Importance sampling LOO, these are <i>negative</i> log-likelihood values). See the Methods (by class) section below for a detailed description of how to specify the inputs for each method.
...	Arguments passed on to the various methods.
r_eff	Vector of relative effective sample size estimates containing one element per observation. The values provided should be the relative effective sample sizes of $1/\exp(\log_ratios)$ (i.e., $1/ratios$). This is related to the relative efficiency of estimating the normalizing term in self-normalizing importance sampling. If <code>r_eff</code> is not provided then the reported (T)IS effective sample sizes and Monte

Carlo error estimates can be over-optimistic. If the posterior draws are (near) independent then `r_eff=1` can be used. `r_eff` has to be a scalar (same value is used for all observations) or a vector with length equal to the number of observations. The default value is 1. See the `relative_eff()` helper function for computing `r_eff`.

cores The number of cores to use for parallelization. This defaults to the option `mc.cores` which can be set for an entire R session by `options(mc.cores = NUMBER)`. The old option `loo.cores` is now deprecated but will be given precedence over `mc.cores` until `loo.cores` is removed in a future release. **As of version 2.0.0 the default is now 1 core if `mc.cores` is not set**, but we recommend using as many (or close to as many) cores as possible.

- Note for Windows 10 users: it is **strongly recommended** to avoid using the `.Rprofile` file to set `mc.cores` (using the `cores` argument or setting `mc.cores` interactively or in a script is fine).

Value

The `tis()` methods return an object of class "tis", which is a named list with the following components:

log_weights Vector or matrix of smoothed (and truncated) but *unnormalized* log weights. To get normalized weights use the `weights()` method provided for objects of class `tis`.

diagnostics A named list containing one vector:

- `pareto_k`: Not used in `tis`, all set to 0.
- `n_eff`: Effective sample size estimates.

Objects of class "tis" also have the following **attributes**:

norm_const_log Vector of precomputed values of `colLogSumExps(log_weights)` that are used internally by the `weights()` method to normalize the log weights.

r_eff If specified, the user's `r_eff` argument.

tail_len Not used for `tis`.

dims Integer vector of length 2 containing `S` (posterior sample size) and `N` (number of observations).

method Method used for importance sampling, here `tis`.

Methods (by class)

- `tis(array)`: An I by C by N array, where I is the number of MCMC iterations per chain, C is the number of chains, and N is the number of data points.
- `tis(matrix)`: An S by N matrix, where S is the size of the posterior sample (with all chains merged) and N is the number of data points.
- `tis(default)`: A vector of length S (posterior sample size).

References

Ionides, Edward L. (2008). Truncated importance sampling. *Journal of Computational and Graphical Statistics* 17(2): 295–311.

See Also

- [psis\(\)](#) for approximate LOO-CV using PSIS.
- [loo\(\)](#) for approximate LOO-CV.
- [pareto-k-diagnostic](#) for PSIS diagnostics.

Examples

```
log_ratios <- -1 * example_loglik_array()
r_eff <- relative_eff(exp(-log_ratios))
tis_result <- tis(log_ratios, r_eff = r_eff)
str(tis_result)

# extract smoothed weights
lw <- weights(tis_result) # default args are log=TRUE, normalize=TRUE
ulw <- weights(tis_result, normalize=FALSE) # unnormalized log-weights

w <- weights(tis_result, log=FALSE) # normalized weights (not log-weights)
uw <- weights(tis_result, log=FALSE, normalize = FALSE) # unnormalized weights
```

update.psis_loo_ss	<i>Update psis_loo_ss objects</i>
--------------------	-----------------------------------

Description

Update psis_loo_ss objects

Usage

```
## S3 method for class 'psis_loo_ss'
update(
  object,
  ...,
  data = NULL,
  draws = NULL,
  observations = NULL,
  r_eff = 1,
  cores = getOption("mc.cores", 1),
  loo_approximation = NULL,
  loo_approximation_draws = NULL,
  llgrad = NULL,
  llhess = NULL
)
```

Arguments

object	A psis_loo_ss object to update.
...	Currently not used.
data, draws	See loo_subsample.function() .
observations	The subsample observations to use. The argument can take four (4) types of arguments: • NULL to use all observations. The algorithm then just uses standard <code>loo()</code> or <code>loo_approximate_posterior()</code>. • A single integer to specify the number of observations to be subsampled. • A vector of integers to provide the indices used to subset the data. <i>These observations need to be subsampled with the same scheme as given by the estimator argument.</i> • A psis_loo_ss object to use the same observations that were used in a previous call to <code>loo_subsample()</code>.
r_eff	Vector of relative effective sample size estimates for the likelihood (<code>exp(log_lik)</code>) of each observation. This is related to the relative efficiency of estimating the normalizing term in self-normalized importance sampling when using posterior draws obtained with MCMC. If MCMC draws are used and <code>r_eff</code> is not provided then the reported PSIS effective sample sizes and Monte Carlo error estimates can be over-optimistic. If the posterior draws are (near) independent then <code>r_eff=1</code> can be used. <code>r_eff</code> has to be a scalar (same value is used for all observations) or a vector with length equal to the number of observations. The default value is 1. See the relative_eff() helper functions for help computing <code>r_eff</code> .
cores	The number of cores to use for parallelization. This defaults to the option <code>mc.cores</code> which can be set for an entire R session by <code>options(mc.cores = NUMBER)</code> . The old option <code>loo.cores</code> is now deprecated but will be given precedence over <code>mc.cores</code> until <code>loo.cores</code> is removed in a future release. As of version 2.0.0 the default is now 1 core if <code>mc.cores</code> is not set , but we recommend using as many (or close to as many) cores as possible. • Note for Windows 10 users: it is strongly recommended to avoid using the <code>.Rprofile</code> file to set <code>mc.cores</code> (using the <code>cores</code> argument or setting <code>mc.cores</code> interactively or in a script is fine).
loo_approximation	What type of approximation of the <code>loo_i</code> 's should be used? The default is "plpd" (the log predictive density using the posterior expectation). There are six different methods implemented to approximate <code>loo_i</code> 's (see the references for more details): • "plpd": uses the lpd based on point estimates (i.e., $p(y_i \hat{\theta})$). • "lpd": uses the lpds (i.e., $p(y_i y)$). • "tis": uses truncated importance sampling to approximate PSIS-LOO. • "waic": uses waic (i.e., $p(y_i y) - p_{waic}$). • "waic_grad_marginal": uses waic approximation using first order delta method and posterior marginal variances to approximate p_{waic} (ie. $p(y_i \hat{\theta}) - p_{waic_grad_marginal}$). Requires gradient of likelihood function.

- "waic_grad": uses waic approximation using first order delta method and posterior covariance to approximate p_{waic} (ie. $p(y_i|\hat{\theta})$ -p_waic_grad). Requires gradient of likelihood function.
- "waic_hess": uses waic approximation using second order delta method and posterior covariance to approximate p_{waic} (ie. $p(y_i|\hat{\theta})$ -p_waic_grad). Requires gradient and Hessian of likelihood function.

As point estimates of $\hat{\theta}$, the posterior expectations of the parameters are used.

loo_approximation_draws

The number of posterior draws used when integrating over the posterior. This is used if loo_approximation is set to "lpd", "waic", or "tis".

llgrad

The gradient of the log-likelihood. This is only used when loo_approximation is "waic_grad", "waic_grad_marginal", or "waic_hess". The default is NULL.

llhess

The Hessian of the log-likelihood. This is only used with loo_approximation = "waic_hess". The default is NULL.

Details

If observations is updated then if a vector of indices or a psis_loo_ss object is supplied the updated object will have exactly the observations indicated by the vector or psis_loo_ss object. If a single integer is supplied, new observations will be sampled to reach the supplied sample size.

Value

A psis_loo_ss object.

waic	<i>Widely applicable information criterion (WAIC)</i>
------	---

Description

The waic() methods can be used to compute WAIC from the pointwise log-likelihood. However, we recommend LOO-CV using PSIS (as implemented by the [loo\(\)](#) function) because PSIS provides useful diagnostics as well as effective sample size and Monte Carlo estimates.

Usage

```
waic(x, ...)

## S3 method for class 'array'
waic(x, ...)

## S3 method for class 'matrix'
waic(x, ...)

## S3 method for class '`function`'
```

```
waic(x, ..., data = NULL, draws = NULL)
```

```
is.waic(x)
```

Arguments

- `x` A log-likelihood array, matrix, or function. The **Methods (by class)** section, below, has detailed descriptions of how to specify the inputs for each method.
- `draws, data, ...` For the function method only. See the **Methods (by class)** section below for details on these arguments.

Value

A named list (of class `c("waic", "loo")`) with components:

- `estimates` A matrix with two columns ("Estimate", "SE") and three rows ("elpd_waic", "p_waic", "waic"). This contains point estimates and standard errors of the expected log pointwise predictive density (elpd_waic), the effective number of parameters (p_waic) and the information criterion waic (which is just $-2 * \text{elpd_waic}$, i.e., converted to deviance scale).
- `pointwise` A matrix with three columns (and number of rows equal to the number of observations) containing the pointwise contributions of each of the above measures (elpd_waic, p_waic, waic).

Methods (by class)

- `waic(array)`: An I by C by N array, where I is the number of MCMC iterations per chain, C is the number of chains, and N is the number of data points.
- `waic(matrix)`: An S by N matrix, where S is the size of the posterior sample (with all chains merged) and N is the number of data points.
- `waic(~function~)`: A function `f()` that takes arguments `data_i` and `draws` and returns a vector containing the log-likelihood for a single observation i evaluated at each posterior draw. The function should be written such that, for each observation i in $1:N$, evaluating

```
f(data_i = data[i,, drop=FALSE], draws = draws)
```

results in a vector of length S (size of posterior sample). The log-likelihood function can also have additional arguments but `data_i` and `draws` are required.

If using the function method then the arguments `data` and `draws` must also be specified in the call to `loo()`:

- `data`: A data frame or matrix containing the data (e.g. observed outcome and predictors) needed to compute the pointwise log-likelihood. For each observation i , the i th row of `data` will be passed to the `data_i` argument of the log-likelihood function.
- `draws`: An object containing the posterior draws for any parameters needed to compute the pointwise log-likelihood. Unlike `data`, which is indexed by observation, for each observation the entire object `draws` will be passed to the `draws` argument of the log-likelihood function.
- The `...` can be used if your log-likelihood function takes additional arguments. These arguments are used like the `draws` argument in that they are recycled for each observation.

References

- Watanabe, S. (2010). Asymptotic equivalence of Bayes cross validation and widely applicable information criterion in singular learning theory. *Journal of Machine Learning Research* **11**, 3571–3594.
- Vehtari, A., Gelman, A., and Gabry, J. (2017). Practical Bayesian model evaluation using leave-one-out cross-validation and WAIC. *Statistics and Computing*. 27(5), 1413–1432. doi:10.1007/s11222-016-9696-4 ([journal version](#), [preprint arXiv:1507.04544](#)).
- Vehtari, A., Simpson, D., Gelman, A., Yao, Y., and Gabry, J. (2024). Pareto smoothed importance sampling. *Journal of Machine Learning Research*, 25(72):1-58. [PDF](#)

See Also

- The **loo** package [vignettes](#) and Vehtari, Gelman, and Gabry (2017) and Vehtari, Simpson, Gelman, Yao, and Gabry (2024) for more details on why we prefer loo() to waic().
- [loo_compare\(\)](#) for comparing models on approximate LOO-CV or WAIC.

Examples

```
### Array and matrix methods
LLarr <- example_loglik_array()
dim(LLarr)

LLmat <- example_loglik_matrix()
dim(LLmat)

waic_arr <- waic(LLarr)
waic_mat <- waic(LLmat)
identical(waic_arr, waic_mat)

## Not run:
log_lik1 <- extract_log_lik(stanfit1)
log_lik2 <- extract_log_lik(stanfit2)
(waic1 <- waic(log_lik1))
(waic2 <- waic(log_lik2))
print(compare(waic1, waic2), digits = 2)

## End(Not run)
```

```
weights.importance_sampling
```

Extract importance sampling weights

Description

Extract importance sampling weights

Usage

```
## S3 method for class 'importance_sampling'  
weights(object, ..., log = TRUE, normalize = TRUE)
```

Arguments

object	An object returned by <code>psis()</code> , <code>tis()</code> , or <code>sis()</code> .
...	Ignored.
log	Should the weights be returned on the log scale? Defaults to TRUE.
normalize	Should the weights be normalized? Defaults to TRUE.

Value

The `weights()` method returns an object with the same dimensions as the `log_weights` component of `object`. The `normalize` and `log` arguments control whether the returned weights are normalized and whether or not to return them on the log scale.

Examples

```
# See the examples at help("psis")
```

Index

ap_psis, 5
attributes, 56, 62, 64

base::round(), 54

compare, 6
constrOptim(), 38
crps, 7

E_loo, 10
E_loo(), 9, 21, 44
ELPD, 33, 34
elpd, 12
elpd_diff, 34
elpd_loo, 22, 33
example_loglik_array, 13
example_loglik_matrix
 (example_loglik_array), 13
extract_log_lik, 14

factor, 19
find_model_names, 15

gpdfit, 15
graphics::text(), 51

importance_sampling, 16
is.kfold(kfold-generic), 18
is.loo(loo), 20
is.psis(psis), 55
is.psis_loo(loo), 20
is.sis(psis), 55
is.tis(psis), 55
is.waic(waic), 67

kfold(kfold-generic), 18
kfold-generic, 18
kfold-helpers, 18
kfold_split_grouped(kfold-helpers), 18
kfold_split_random(kfold-helpers), 18
kfold_split_stratified(kfold-helpers),
 18
Kline (loo-datasets), 25

loo, 20
loo(), 3, 6, 31–33, 36–38, 42, 43, 47–50,
 52–54, 57–60, 62, 65, 67
loo-datasets, 25
loo-glossary, 26
loo-package, 3
loo.array(loo), 20
loo.function(loo), 20
loo.function(), 58
loo.matrix(loo), 20
loo_approximate_posterior, 30
loo_approximate_posterior(), 46, 53
loo_compare, 33
loo_compare(), 6, 18, 23, 33, 49, 69
loo_crps(crps), 7
loo_i(loo), 20
loo_model_weights, 35
loo_model_weights(), 3
loo_moment_match, 40
loo_moment_match(), 43, 52, 53
loo_moment_match_split, 42
loo_moment_match_split(), 42
loo_predictive_metric, 44
loo_scrps(crps), 7
loo_subsample, 45
loo_subsample(), 53
loo_subsample.function(), 66

mcse_elpd_loo, 22
mcse_loo(pareto-k-diagnostic), 50
milk(loo-datasets), 25

nobs.psis_loo_ss, 49

obs_idx, 49

p_loo, 22

Pareto k, [10](#), [11](#)
 pareto-k-diagnostic, [16](#), [22](#), [23](#), [50](#), [55–58](#),
 [62](#), [65](#)
 pareto_k_ids (pareto-k-diagnostic), [50](#)
 pareto_k_influence_values
 (pareto-k-diagnostic), [50](#)
 pareto_k_table (pareto-k-diagnostic), [50](#)
 pareto_k_values (pareto-k-diagnostic),
 [50](#)
 plot(), [54](#)
 plot.loo (pareto-k-diagnostic), [50](#)
 plot.psis (pareto-k-diagnostic), [50](#)
 plot.psis_loo (pareto-k-diagnostic), [50](#)
 plot.psis_loo(), [54](#)
 pointwise, [53](#)
 print.compare.loo (loo_compare), [33](#)
 print.compare.loo_ss (loo_compare), [33](#)
 print.importance_sampling (print.loo),
 [54](#)
 print.importance_sampling_loo
 (print.loo), [54](#)
 print.kfold (print.loo), [54](#)
 print.loo, [54](#)
 print.psis (print.loo), [54](#)
 print.psis_loo (print.loo), [54](#)
 print.psis_loo_ap (print.loo), [54](#)
 print.waic (print.loo), [54](#)
 pseudobma_weights (loo_model_weights),
 [35](#)
 PSIS, [10](#), [20](#)
 psis, [55](#)
 psis(), [3](#), [9](#), [11](#), [16](#), [21–23](#), [33](#), [37](#), [44](#), [49](#), [50](#),
 [52–54](#), [57](#), [59](#), [62](#), [65](#), [70](#)
 psis_n_eff_values
 (pareto-k-diagnostic), [50](#)
 psislw, [57](#)

 relative_eff, [59](#)
 relative_eff(), [17](#), [21](#), [37](#), [38](#), [47](#), [56](#), [61](#),
 [64](#), [66](#)

 scrps (crps), [7](#)
 se_diff, [34](#)
 sis, [61](#)
 sis(), [70](#)
 stacking_weights (loo_model_weights), [35](#)
 stats::constrOptim(), [36](#)
 stats::optim(), [36](#)

 tis, [63](#)
 tis(), [70](#)

 update.psis_loo_ss, [65](#)

 voice (loo-datasets), [25](#)
 voice_loo (loo-datasets), [25](#)

 waic, [67](#)
 waic(), [6](#), [54](#)
 weights(), [56](#), [62](#), [64](#)
 weights.importance_sampling, [69](#)