

Package ‘FDRsamplesize2’

July 21, 2025

Type Package

Title Computing Power and Sample Size for the False Discovery Rate in Multiple Applications

Version 0.2.0

Maintainer Yonghui Ni <Yonghui.Ni@STJUDE.ORG>

Description Defines a collection of functions to compute average power and sample size for studies that use the false discovery rate as the final measure of statistical significance. A three-rectangle approximation method of a p-value histogram is proposed to derive a formula to compute the statistical power for analyses that involve the FDR. The methodology paper of this package is under review.

Encoding UTF-8

License MIT + file LICENSE

RoxygenNote 7.2.3

Suggests testthat (>= 3.0.0)

Config/testthat/edition 3

NeedsCompilation no

Author Yonghui Ni [aut, cre],
Stanley Pounds [aut]

Repository CRAN

Date/Publication 2024-02-27 03:50:03 UTC

Contents

alpha.power.fdr	2
average.power.coxph	3
average.power.fisher	4
average.power.hart	5
average.power.li	6
average.power.oneway	7
average.power.ranksum	8
average.power.signrank	8

average.power.signtest	9
average.power.t.test	10
average.power.tcorr	11
average.power.twoprop	11
fdr.avepow	12
fdr.power.alpha	13
find.sample.size	14
n.fdr.coxph	16
n.fdr.fisher	17
n.fdr.negbin	18
n.fdr.oneway	19
n.fdr.poisson	20
n.fdr.ranksum	21
n.fdr.signrank	22
n.fdr.signtest	23
n.fdr.tcorr	24
n.fdr.ttest	25
n.fdr.twoprop	27
power.cox	28
power.fisher	29
power.hart	29
power.li	30
power.oneway	31
power.ranksum	32
power.signrank	33
power.signtest	34
power.tcorr	35
Index	36

alpha.power.fdr	<i>Compute p-value threshold for given the proportion π_0 of tests with a true null</i>
-----------------	--

Description

Given the proportion π_0 of tests with a true null, find the p-value threshold that results in a desired FDR and average power.

Usage

```
alpha.power.fdr(fdr, pwr, pi0, method = "HH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
pi0	the proportion of tests with a true null hypothesis
method	method to estimate proportion pi0 of tests with true null, including: "HH" (p-value histogram height) , "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Details

To get the fixed p-value threshold for multiple testing procedure, 4 approximation methods are provided, they are Benjamini & Hochberg procedure (1995), Jung's formula (2005), method of using p-value histogram height (HH) and method of using p-value histogram mean (HM). For last two methods' details, see Ni Y, Onar-Thomas A, Pounds S. "Computing Power and Sample Size for the False Discovery Rate in Multiple Applications"

Value

The fixed p-value threshold for multiple testing procedure

References

- Pounds S and Cheng C, "Sample size determination for the false discovery rate." *Bioinformatics* 21.23 (2005): 4263-4271.
- Gadbury GL, et al. (2004) Power and sample size estimation in high dimensional biology. *Statistical Methods in Medical Research* 13(4):325-38.
- Jung,Sin-Ho."Sample size for FDR-control in microarray data analysis." *Bioinformatics* 21.14 (2005): 3097-3104.
- Ni Y, Seffernick A, Onar-Thomas A, Pounds S. "Computing Power and Sample Size for the False Discovery Rate in Multiple Applications", Manuscript.

Examples

```
alpha.power.fdr(fdr = 0.1, pwr = 0.9, pi0=0.9, method = "HH")
```

average.power.coxph *Compute the average power of many Cox regression models*

Description

Compute the average power of many Cox regression models for a given number of events, p-value threshold, vector of effect sizes (log hazard ratio),and variance of predictor variables

Usage

```
average.power.coxph(n, alpha, logHR, v)
```

Arguments

n	number of events (scalar)
alpha	p-value threshold (scalar)
logHR	log hazard ratio (vector)
v	variance of predictor variable (vector)

Value

Average power estimate for multiple testing procedure

References

Hsieh, FY and Lavori, Philip W (2000) Sample-size calculations for the Cox proportional hazards regression model with non-binary covariates. *Controlled Clinical Trials* 21(6):552-560.

See Also

[power.cox](#) for more details about power calculation of single-predictor Cox regression model. The power calculation is based on asymptotic normal approximation.

Examples

```
logHR = log(rep(c(1, 2), c(900, 100)));
v = rep(1, 1000);
average.power.coxph(n = 50, alpha = 0.05, logHR = logHR, v = v)
```

average.power.fisher *Compute average power of many Fisher's exact tests*

Description

Compute average power of many Fisher's exact tests

Usage

```
average.power.fisher(p1, p2, n, alpha, alternative)
```

Arguments

p1	probability in one group (vector)
p2	probability in other group (vector)
n	per-group sample size
alpha	p-value threshold
alternative	one- or two-sided test

Value

Average power estimate for multiple testing procedure

See Also

[power.fisher](#) for more details about power calculation of Fisher's exact test

Examples

```
set.seed(1234);
p1 = sample(seq(0,0.5,0.1),5,replace = TRUE);
p2 = sample(seq(0.5,1,0.1),5,replace = TRUE);
average.power.fisher(p1 = p1,p2 = p2,n = 20,alpha = 0.05,alternative = "two.sided")
```

average.power.hart	<i>Compute average power for RNA-seq experiments assuming Negative Binomial distribution</i>
--------------------	--

Description

Compute average power for RNA-seq experiments assuming Negative Binomial distribution

Usage

```
average.power.hart(n, alpha, log.fc, mu, sig)
```

Arguments

n	per-group sample size (scalar)
alpha	p-value threshold (scalar)
log.fc	log fold-change (vector), usual null hypothesis is log.fc=0
mu	read depth per gene (vector, same length as log.fc)
sig	coefficient of variation (CV) per gene (vector, same length as log.fc)

Details

The power function is based on equation (1) of Hart et al (2013). It assumes a Negative Binomial model for RNA-seq read counts and equal sample size per group.

Value

Average power estimate for multiple testing procedure

References

SN Hart, TM Therneau, Y Zhang, GA Poland, and J-P Kocher (2013). Calculating Sample Size Estimates for RNA Sequencing Data. *Journal of Computational Biology* 20: 970-978.

See Also

[power.hart](#) for more details about power calculation of data under Negative Binomial distribution. The power calculation is based on asymptotic normal approximation.

Examples

```
logFC = log(rep(c(1,2),c(900,100)));
mu = rep(5,1000);
sig = rep(0.6,1000);
average.power.hart(n = 50, alpha = 0.05, log.fc = logFC, mu = mu, sig = sig)
```

average.power.li	<i>Compute average power for RNA-Seq experiments assuming Poisson distribution</i>
------------------	--

Description

Use the formula of Li et al (2013) to compute power for comparing RNA-seq expression across two groups assuming the Poisson distribution.

Usage

```
average.power.li(n, alpha, rho, mu0, w, type)
```

Arguments

n	per-group sample size
alpha	p-value threshold (scalar)
rho	fold-change, usual null hypothesis is that rho=1 (vector)
mu0	average count in control group (vector)
w	ratio of the total number of reads mapped between the two groups (scalar or vector)
type	type of test: "w" for Wald, "s" for score, "lw" for log-transformed Wald, "ls" for log-transformed score

Details

This function computes the average power for a series of two-sided tests defined by the input parameters. The power is based on the sample size formulas in equations (10-13) of Li et al (2013). Also, note that the null.effect is set to 1 in the examples because the usual null hypothesis is that the fold-change = 1.

Value

Average power estimate for multiple testing procedure

References

C-I Li, P-F Su, Y Guo, and Y Shyr (2013). Sample size calculation for differential expression analysis of RNA-seq data under Poisson distribution. Int J Comput Biol Drug Des 6(4).<doi:10.1504/IJCBDD.2013.056830>

See Also

[power.li](#) for more details about power calculation of data under Poisson distribution

Examples

```
rho = rep(c(1,1.25),c(900,100));
mu0 = rep(5,1000);
w = rep(0.5,1000);
average.power.li(n = 50, alpha = 0.05, rho = rho, mu0 = mu0, w = w, type = "w")
```

average.power.oneway *Compute average power of many one-way ANOVA tests*

Description

Compute average power of many one-way ANOVA tests

Usage

```
average.power.oneway(n, alpha, theta, k)
```

Arguments

n	per-group sample size (scalar)
alpha	p-value threshold (scalar)
theta	sum of ((group mean - overall mean)/stdev)^2 across all groups for each hypothesis test(vector)
k	the number of groups to be compared

Value

Average power estimate for multiple testing procedure

See Also

[power.oneway](#) for more details about power calculation of one-way ANOVA

Examples

```
theta=rep(c(2,0),c(100,900));
average.power.oneway(n = 50, alpha = 0.05, theta = theta, k = 2)
```

average.power.ranksum *Compute average power of rank-sum tests*

Description

Compute average power of rank-sum tests

Usage

```
average.power.ranksum(n, alpha, p)
```

Arguments

n	sample size (scalar)
alpha	p-value threshold (scalar)
p	$\Pr(Y>X)$, as in Noether (JASA 1987)

Value

Average power estimate for multiple testing procedure

See Also

[power.ranksum](#) for more details about power calculation of rank-sum test. The power calculation is based on asymptotic normal approximation.

Examples

```
p = rep(c(0.8,0.5),c(100,900));
average.power.ranksum(n = 50, alpha = 0.05, p=p)
```

average.power.signrank
Compute average power of many signed-rank tests

Description

Compute average power of many signed-rank tests

Usage

```
average.power.signrank(n, alpha, p1, p2)
```

Arguments

n	sample size (scalar)
alpha	p-value threshold (scalar)
p1	$\Pr(X>0)$, as in Noether (JASA 1987)
p2	$\Pr(X+X'>0)$, as in Noether (JASA 1987)

Value

Average power estimate for multiple testing procedure

See Also

[power.signrank](#) for more details about power calculation of signed-rank test. The power calculation is based on asymptotic normal approximation.

Examples

```
p1 = rep(c(0.8,0.5),c(100,900));
p2 = rep(c(0.8,0.5),c(100,900));
average.power.signrank(n = 50, alpha = 0.05, p1 = p1, p2 = p2)
```

average.power.signtest

Compute average power of many sign tests

Description

Compute average power of many sign tests

Usage

```
average.power.signtest(n, alpha, p)
```

Arguments

n	sample size (scalar)
alpha	p-value threshold (scalar)
p	$\Pr(Y>X)$, as in Noether (JASA 1987)

Value

Average power estimate for multiple testing procedure

See Also

[power.signtest](#) for more details about power calculation of sign test. The power calculation is based on asymptotic normal approximation.

Examples

```
p = rep(c(0.8,0.5),c(100,900));
average.power.signtest(n = 50, alpha = 0.05, p=p)
```

average.power.t.test *Compute average power of many t-tests*

Description

Compute average power of many t-tests; Uses classical power formula for t-test; Assumes equal variance and sample size

Usage

```
average.power.t.test(
  n,
  alpha,
  delta,
  sigma = 1,
  type = "two.sample",
  alternative = "two.sided"
)
```

Arguments

n	per-group sample size (scalar)
alpha	p-value threshold (scalar)
delta	difference of population means (vector)
sigma	standard deviation (vector or scalar, default=1)
type	type of t-test: "two.sample", "one.sample"
alternative	one- or two-sided test

Value

Average power estimate for multiple testing procedure

Examples

```
d = rep(c(2,0),c(100,900));
average.power.t.test(n = 20, alpha = 0.05,delta = d)
```

average.power.tcorr	<i>Compute average power of many t-tests for non-zero correlation</i>
---------------------	---

Description

Compute average power of many t-tests for non-zero correlation

Usage

```
average.power.tcorr(n, alpha, rho)
```

Arguments

n	sample size (scalar)
alpha	p-value threshold (scalar)
rho	population correlation coefficient (vector)

Details

For many applications, the null.effect is $\rho = 0$

Value

Average power estimate for multiple testing procedure

See Also

[power.tcorr](#) for more details about power calculation of t-test for non-zero correlation

Examples

```
rho = rep(c(0.3,0),c(100,900));  
average.power.tcorr(n = 50, alpha = 0.05, rho = rho)
```

average.power.twoprop	<i>Computer average power of many two proportion z-tests</i>
-----------------------	--

Description

Computer average power of many two proportion z-tests. The power calculation of two proportion z-test is based on asymptotic normal approximation.

Usage

```
average.power.twoprop(n, alpha, p1, p2, alternative)
```

Arguments

n	per-group sample size (scalar)
alpha	p-value threshold (scalar)
p1	probability in one group (vector)
p2	probability in other group (vector)
alternative	one- or two-sided test

Value

Average power estimate for multiple testing procedure

Examples

```
set.seed(1234);
p1 = sample(seq(0,0.5,0.1),40,replace = TRUE);
p2 = sample(seq(0.5,1,0.1),40,replace = TRUE);
average.power.twoprop(n = 30, alpha = 0.05, p1 = p1,p2 = p2,alternative="two.sided")
```

fdr.avepow	<i>Compute FDR and average power for a given sample size and effect size vector</i>
------------	---

Description

For a given fixed sample size and effect size vector, compute FDR and average power as a function of the p-value threshold alpha.

Usage

```
fdr.avepow(n, avepow.func, null.hypo, alpha = 1:100/1000, method = "BH", ...)
```

Arguments

n	sample size
avepow.func	function to compute average power
null.hypo	string to evaluate null hypothesis
alpha	p-value threshold(s) to consider
method	method to estimate proportion π_0 of tests with a true null hypothesis, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)
...	additional arguments, including effect size vector for average power function

Value

A list with the following components:

<code>n</code>	input sample size
<code>avepow.func</code>	average power function
<code>null.hypo</code>	null hypothesis string
<code>pi0</code>	computed value of π_0
<code>method</code>	method to estimate proportion π_0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)
<code>other.args</code>	additional arguments
<code>res.tbl</code>	table of alpha, fdr, and average power

References

- Pounds S and Cheng C, "Sample size determination for the false discovery rate." *Bioinformatics* 21.23 (2005): 4263-4271.
- Gadbury GL, et al. (2004) Power and sample size estimation in high dimensional biology. *Statistical Methods in Medical Research* 13(4):325-38.
- Jung,Sin-Ho."Sample size for FDR-control in microarray data analysis." *Bioinformatics* 21.14 (2005): 3097-3104.
- Ni Y, Seffernick A, Onar-Thomas A, Pounds S. "Computing Power and Sample Size for the False Discovery Rate in Multiple Applications", Manuscript.

Examples

```
n = 50; # number of events
logHR = rep(c(0,0.5),c(950,50));
v = rep(1,length(logHR)); # variance of predictor variable (vector)
res = fdr.avepow(n,average.power.coxph,"logHR==0",logHR=logHR,v=v);
res$pi0;
head(res$res.tbl)
```

<code>fdr.power.alpha</code>	<i>Compute FDR for given p-value threshold, average power and proportion of tests with a true null</i>
------------------------------	--

Description

Compute the FDR for given values of the p-value threshold alpha, average power, and proportion π_0 of tests with a true null hypothesis.

Usage

```
fdr.power.alpha(alpha, pwr, pi0, method = "HH")
```

Arguments

<code>alpha</code>	p-value threshold (vector)
<code>pwr</code>	average power
<code>pi0</code>	actual proportion of tests with a true null hypothesis
<code>method</code>	method to estimate proportion <code>pi0</code> of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

FDR

References

- Pounds S and Cheng C, "Sample size determination for the false discovery rate." *Bioinformatics* 21.23 (2005): 4263-4271.
- Gadbury GL, et al. (2004) Power and sample size estimation in high dimensional biology. *Statistical Methods in Medical Research* 13(4):325-38.
- Jung,Sin-Ho."Sample size for FDR-control in microarray data analysis." *Bioinformatics* 21.14 (2005): 3097-3104.
- Ni Y, Seffernick A, Onar-Thomas A, Pounds S. "Computing Power and Sample Size for the False Discovery Rate in Multiple Applications", Manuscript.

Examples

```
alpha = 1:100/1000;
pwr = rep(0.8,length(alpha));
pi0 = 0.95;
fdr.power.alpha(alpha,pwr,pi0,method="HH")
```

<code>find.sample.size</code>	<i>Determines the sample size needed to achieve the desired FDR and average power</i>
-------------------------------	---

Description

Determines the sample size needed to achieve the desired FDR and average power by given the proportion of true null hypothesis.

Usage

```
find.sample.size(alpha, pwr, avepow.func, n0 = 3, n1 = 6, max.its = 50, ...)
```

Arguments

<code>alpha</code>	the fixed p-value threshold (scalar numeric)
<code>pwr</code>	desired average power (scalar numeric)
<code>avepow.func</code>	an R function to compute average power
<code>n0</code>	lower limit for initial sample size range
<code>n1</code>	upper limit for initial sample size range
<code>max.its</code>	maximum number of iterations
<code>...</code>	additional arguments to average power function

Value

A list with the following components:

<code>n</code>	a sample size estimate
<code>computed.avepow</code>	average power
<code>desired.avepow</code>	desired average power
<code>alpha</code>	fixed p-value threshold for multiple testing procedure
<code>n.its</code>	number of iteration
<code>max.its</code>	maximum number of iteration, default is 50
<code>n0</code>	lower limit for initial sample size range
<code>n1</code>	upper limit for initial sample size range

Note

For the test with power calculation based on asymptotic normal approximation, we suggest checking `FDRsample.size2` calculation by simulation.

Examples

```
#Here, calculating the sample size for the study involving many sign tests
average.power.signtest;
p.adj = 0.001;
p = rep(c(0.8,0.5), c(100,9900));
find.sample.size(alpha = p.adj, pwr = 0.8, avepow.func = average.power.signtest, p = p)
```

n.fdr.coxph	<i>Sample size calculation for the Cox proportional hazards regression model</i>
-------------	--

Description

Find number of events needed to have a desired false discovery rate and average power for a large number of Cox regression models with non-binary covariates.

Usage

```
n.fdr.coxph(fdr, pwr, logHR, v, pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
logHR	log hazard ratio (vector)
v	variance of predictor variable (vector)
pi0.hat	method to estimate proportion π_0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

n	number of events estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis
alpha	fixed p-value threshold for multiple testing procedure
n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

Note

For the test with power calculation based on asymptotic normal approximation, we suggest checking `FDRsampleSize2` calculation by simulation.

References

Hsieh, FY and Lavori, Philip W (2000) Sample-size calculations for the Cox proportional hazards regression model with non-binary covariates. *Controlled Clinical Trials* 21(6):552-560.

Examples

```
log.HR=log(rep(c(1,2),c(900,100)))
v=rep(1,1000)
n.fdr.coxph(fdr=0.1, pwr=0.8,logHR=log.HR, v=v, pi0.hat="BH")
```

n.fdr.fisher

Sample size calculation for Fisher's Exact tests

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of Fisher's exact tests.

Usage

```
n.fdr.fisher(fdr, pwr, p1, p2, alternative = "two.sided", pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
p1	probability in one group (vector)
p2	probability in other group (vector)
alternative	one- or two-sided test
pi0.hat	method to estimate proportion π_0 of tests with true null, including: "HH" (p-value histogram height) , "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

n	per-group sample size estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis
alpha	fixed p-value threshold for multiple testing procedure

n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

Examples

```
set.seed(1234);
p1 = sample(seq(0,0.5,0.1),10,replace = TRUE);
p2 = sample(seq(0.5,1,0.1),10,replace = TRUE);
n.fdr.fisher(fdr = 0.1, pwr = 0.8, p1 = p1, p2 = p2, alternative = "two.sided", pi0.hat = "BH")
```

n.fdr.negbin	<i>Sample size calculation for Negative Binomial data</i>
--------------	---

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of Negative Binomial comparisons.

Usage

```
n.fdr.negbin(fdr, pwr, log.fc, mu, sig, pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
log.fc	log fold-change (vector), usual null hypothesis is log.fc=0
mu	read depth per gene (vector, same length as log.fc)
sig	coefficient of variation (CV) per gene (vector, same length as log.fc)
pi0.hat	method to estimate proportion pi0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

n	per-group sample size estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis

alpha	fixed p-value threshold for multiple testing procedure
n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

Note

For the test with power calculation based on asymptotic normal approximation, we suggest checking FDRsampleSize2 calculation by simulation.

References

SN Hart, TM Therneau, Y Zhang, GA Poland, and J-P Kocher (2013). Calculating Sample Size Estimates for RNA Sequencing Data. *Journal of Computational Biology* 20: 970-978.

Examples

```
logFC = log(rep(c(1,2),c(900,100)));
mu = rep(5,1000);
sig = rep(0.6,1000);
n.fdr.negbin(fdr = 0.1, pwr = 0.8, log.fc = logFC, mu = mu, sig = sig, pi0.hat = "BH")
```

n.fdr.oneway	<i>Sample size calculation for one-way ANOVA</i>
--------------	--

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of one-way ANOVA tests.

Usage

```
n.fdr.oneway(fdr, pwr, theta, k, pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
theta	sum of ((group mean - overall mean)/stdev)^2 across all groups for each hypothesis test (vector)
k	the number of groups to be compared
pi0.hat	method to estimate proportion pi0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

n	per-group sample size estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis
alpha	fixed p-value threshold for multiple testing procedure
n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

Examples

```
theta=rep(c(2,0),c(100,900));
n.fdr.oneway(fdr = 0.1, pwr = 0.8, theta = theta, k = 2, pi0.hat = "BH")
```

n.fdr.poisson	<i>Sample size calculation for Poisson data</i>
---------------	---

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of two-group comparisons under Poisson distribution.

Usage

```
n.fdr.poisson(fdr, pwr, rho, mu0, w, type, pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
rho	fold-change, usual null hypothesis is that rho=1 (vector)
mu0	average count in control group (vector)
w	ratio of the total number of reads mapped between the two groups
type	type of test: "w" for Wald, "s" for score, "lw" for log-transformed Wald, "ls" for log-transformed score.
pi0.hat	method to estimate proportion pi0 of tests with true null, including: "HH" (p-value histogram height) , "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

n	per-group sample size estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis
alpha	fixed p-value threshold for multiple testing procedure
n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

References

C-I Li, P-F Su, Y Guo, and Y Shyr (2013). Sample size calculation for differential expression analysis of RNA-seq data under Poisson distribution. Int J Comput Biol Drug Des 6(4).<doi:10.1504/IJCBDD.2013.056830>

Examples

```
rho = rep(c(1,1.25),c(900,100));
mu0 = rep(5,1000);
w = rep(0.5,1000);
n.fdr.poisson(fdr = 0.1, pwr = 0.8, rho = rho, mu0 = mu0, w = w, type = "w", pi0.hat = "BH")
```

n.fdr.ranksum	<i>Sample size calculation for rank-sum tests</i>
---------------	---

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of rank-sum tests.

Usage

```
n.fdr.ranksum(fdr, pwr, p, pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
p	$\Pr(Y>X)$, as in Noether (JASA 1987)
pi0.hat	method to estimate proportion π_0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

n	sample size estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis
alpha	fixed p-value threshold for multiple testing procedure
n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

References

Noether, Gottfried E (1987) Sample size determination for some common nonparametric tests. Journal of the American Statistical Association, 82:645-647.

Examples

```
p = rep(c(0.8,0.5),c(100,900));
n.fdr.ranksum(fdr = 0.1, pwr = 0.8, p = p, pi0.hat = "BH")
```

n.fdr.signrank	<i>Sample size calculation for signed-rank tests</i>
----------------	--

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of signed-rank tests.

Usage

```
n.fdr.signrank(fdr, pwr, p1, p2, pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
p1	$\Pr(X > 0)$, as in Noether (JASA 1987)
p2	$\Pr(X + X' > 0)$, as in Noether (JASA 1987)
pi0.hat	method to estimate proportion π_0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

n	sample size estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis
alpha	fixed p-value threshold for multiple testing procedure
n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

References

Noether, Gottfried E (1987) Sample size determination for some common nonparametric tests. Journal of the American Statistical Association, 82:645-647.

Examples

```
p1 = rep(c(0.8,0.5),c(100,900));
p2 = rep(c(0.8,0.5),c(100,900));
n.fdr.signrank(fdr = 0.1, pwr = 0.8, p1 = p1, p2 = p2, pi0.hat = "BH")
```

n.fdr.signtest	<i>Sample size calculation for sign tests</i>
----------------	---

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of sign tests.

Usage

```
n.fdr.signtest(fdr, pwr, p, pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
p	$\Pr(X>0)$, as in Noether (JASA 1987)
pi0.hat	method to estimate proportion π_0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

<code>n</code>	sample size estimate
<code>computed.avepow</code>	average power
<code>desired.avepow</code>	desired average power
<code>desired.fdr</code>	desired FDR
<code>input.pi0</code>	proportion of tests with a true null hypothesis
<code>alpha</code>	fixed p-value threshold for multiple testing procedure
<code>n.its</code>	number of iteration
<code>max.its</code>	maximum number of iteration, default is 50
<code>n0</code>	lower limit for initial sample size range
<code>n1</code>	upper limit for initial sample size range

Note

For the test with power calculation based on asymptotic normal approximation, we suggest checking `FDRsamplesize2` calculation by simulation.

References

Noether, Gottfried E (1987) Sample size determination for some common nonparametric tests. Journal of the American Statistical Association, 82:645-647.

Examples

```
p = rep(c(0.8, 0.5), c(100, 900));
n.fdr.signtest(fdr = 0.1, pwr = 0.8, p = p, pi0.hat = "BH")
```

<code>n.fdr.tcorr</code>	<i>Sample size calculation for t-tests for non-zero correlation</i>
--------------------------	---

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of t-tests for non-zero correlation.

Usage

```
n.fdr.tcorr(fdr, pwr, rho, pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
rho	population correlation coefficient (vector)
pi0.hat	method to estimate proportion π_0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

n	sample size estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis
alpha	fixed p-value threshold for multiple testing procedure
n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

Examples

```
rho = rep(c(0.3,0),c(100,900));
n.fdr.tcorr(fdr = 0.1, pwr = 0.8, rho = rho, pi0.hat="BH")
```

n.fdr.ttest

Sample size calculation for t-tests

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of t-tests.

Usage

```
n.fdr.ttest(
  fdr,
  pwr,
  delta,
  sigma = 1,
  type = "two.sample",
  pi0.hat = "BH",
  alternative = "two.sided"
)
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
delta	difference of population means (vector)
sigma	standard deviation (vector or scalar)
type	type of t-test
pi0.hat	method to estimate proportion π_0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)
alternative	one- or two-sided test

Value

A list with the following components:

n	sample size (per group) estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis
alpha	fixed p-value threshold for multiple testing procedure
n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

Examples

```
d = rep(c(2,0),c(100,900));
n.fdr.ttest(fdr = 0.1, pwr = 0.8, delta = d)
```

n.fdr.twoprop

Sample size calculation for comparing two proportions

Description

Find the sample size needed to have a desired false discovery rate and average power for a large number of two-group comparisons using the two proportion z-test.

Usage

```
n.fdr.twoprop(fdr, pwr, p1, p2, alternative = "two.sided", pi0.hat = "BH")
```

Arguments

fdr	desired FDR (scalar numeric)
pwr	desired average power (scalar numeric)
p1	probability in one group (vector)
p2	probability in other group (vector)
alternative	one- or two-sided test
pi0.hat	method to estimate proportion π_0 of tests with true null, including: "HH" (p-value histogram height), "HM" (p-value histogram mean), "BH" (Benjamini & Hochberg 1995), "Jung" (Jung 2005)

Value

A list with the following components:

n	per-group sample size estimate
computed.avepow	average power
desired.avepow	desired average power
desired.fdr	desired FDR
input.pi0	proportion of tests with a true null hypothesis
alpha	fixed p-value threshold for multiple testing procedure
n.its	number of iteration
max.its	maximum number of iteration, default is 50
n0	lower limit for initial sample size range
n1	upper limit for initial sample size range

Note

For the test with power calculation based on asymptotic normal approximation, we suggest checking `FDRsampleSize2` calculation by simulation.

Examples

```
set.seed(1234);
p1 = sample(seq(0,0.5,0.1),40,replace = TRUE);
p2 = sample(seq(0.5,1,0.1),40,replace = TRUE);
n.fdr.twoprop(fdr = 0.1, pwr = 0.8, p1 = p1, p2 = p2, alternative = "two.sided", pi0.hat = "BH")
```

power.cox

*Compute the power of a single-predictor Cox regression model***Description**

Use the formula of Hsieh and Lavori (2000) to compute the power of a single-predictor Cox model, which is based on asymptotic normal approximation.

Usage

```
power.cox(n, alpha, logHR, v)
```

Arguments

n	number of events (scalar)
alpha	p-value threshold (scalar)
logHR	log hazard ratio (vector)
v	variance of predictor variable (vector)

Value

Vector of power estimates for two-sided test

References

Hsieh, FY and Lavori, Philip W (2000) Sample-size calculations for the Cox proportional hazards regression model with non-binary covariates. *Controlled Clinical Trials* 21(6):552-560.

Examples

```
logHR = log(rep(c(1, 2),c(900, 100)));
v = rep(1, 1000);
res = power.cox(n = 50,alpha = 0.05,logHR = logHR, v = v)
```

power.fisher	<i>Compute power for Fisher's exact test</i>
--------------	--

Description

Compute power for Fisher's exact test

Usage

```
power.fisher(p1, p2, n, alpha, alternative)
```

Arguments

p1	probability in one group (scalar)
p2	probability in other group (scalar)
n	per-group sample size (scalar)
alpha	p-value threshold (scalar)
alternative	one- or two-sided test, must be one of "greater", "less", or "two.sided"

Value

Power estimate for one- or two-sided tests

Examples

```
power.fisher(p1 = 0.5, p2 = 0.9, n=20, alpha = 0.05, alternative = 'two.sided')
```

power.hart	<i>Compute power for RNA-seq experiments assuming Negative Binomial distribution</i>
------------	--

Description

Use the formula of Hart et al (2013) to compute power for comparing RNA-seq expression across two groups assuming a Negative Binomial distribution. The power calculation is based on asymptotic normal approximation.

Usage

```
power.hart(n, alpha, log.fc, mu, sig)
```

Arguments

n	per-group sample size (scalar)
alpha	p-value threshold (scalar)
log.fc	log fold-change (vector), usual null hypothesis is log.fc=0
mu	read depth per gene (vector, same length as log.fc)
sig	coefficient of variation (CV) per gene (vector, same length as log.fc)

Details

This function is based on equation (1) of Hart et al (2013). It assumes a Negative Binomial model for RNA-seq read counts and equal sample size per group.

Value

Vector of power estimates for the set of two-sided tests

References

SN Hart, TM Therneau, Y Zhang, GA Poland, and J-P Kocher (2013). Calculating Sample Size Estimates for RNA Sequencing Data. *Journal of Computational Biology* 20: 970-978.

Examples

```
n.hart = 2*(qnorm(0.975)+qnorm(0.9))^2*(1/20+0.6^2)/(log(2)^2) # Equation (6) of Hart et al
power.hart(n.hart,0.05,log(2),20,0.6) # Recapitulate 90% power
```

power.li	<i>Compute power for RNA-Seq experiments assuming Poisson distribution</i>
----------	--

Description

Use the formula of Li et al (2013) to compute power for comparing RNA-seq expression across two groups assuming the Poisson distribution

Usage

```
power.li(n, alpha, rho, mu0, w, type)
```

Arguments

n	per-group sample size
alpha	p-value threshold (scalar)
rho	fold-change, usual null hypothesis is that rho=1 (vector)
mu0	average count in control group
w	ratio of the total number of reads mapped between the two groups
type	type of test: "w" for Wald, "s" for score, "lw" for log-transformed Wald, "ls" for log-transformed score

Details

This function computes the power for each of a series of two-sided tests defined by the input parameters. The power is based on the sample size formulas in equations (10-13) of Li et al (2013). Also, note that the null.effect is set to 1 in the examples because the usual null hypothesis is that the fold-change = 1.

Value

Vector of power estimates for two-sided tests

References

C-I Li, P-F Su, Y Guo, and Y Shyr (2013). Sample size calculation for differential expression analysis of RNA-seq data under Poisson distribution. *Int J Comput Biol Drug Des* 6(4). <doi:10.1504/IJCBDD.2013.056830>

Examples

```
power.li(n = 88, alpha = 0.05, rho = 1.25, mu0 = 5, w = 0.5, type = "w")
# recapitulate 80% power in Table 1 of Li et al (2013)
```

power.oneway

Compute power of one-way ANOVA

Description

Compute power of one-way ANOVA; Uses classical power formula for ANOVA; Assumes equal variance and sample size

Usage

```
power.oneway(n, alpha, theta, k = 2)
```

Arguments

n	per-group sample size (scalar)
alpha	p-value threshold (scalar)
theta	sum of ((group mean - overall mean)/stdev)^2 across all groups for each hypothesis test(vector)
k	the number of groups to be compared, default k=2

Details

For many applications, the null effect is zero for the parameter theta described above

Value

Vector of power estimates for test of equal means

Examples

```
theta=rep(c(2,0),c(100,900));
res = power.oneway(n = 50, alpha = 0.05, theta = theta, k = 2)
```

power.ranksum	<i>Compute power of the rank-sum test</i>
---------------	---

Description

Compute power of rank-sum test; Uses formula of Noether (JASA 1987), which is based on asymptotic normal approximation.

Usage

```
power.ranksum(n, alpha, p)
```

Arguments

n	sample size (scalar)
alpha	p-value threshold (scalar)
p	$\Pr(Y>X)$, as in Noether (JASA 1987)

Details

In most applications, the null effect size will be designated by $p = 0.5$

Value

Vector of power estimates for two-sided tests

References

Noether, Gottfried E (1987) Sample size determination for some common nonparametric tests. Journal of the American Statistical Association, 82:645-647.

Examples

```
p = rep(c(0.8,0.5),c(100,900))
res = power.ranksum(n = 50, alpha = 0.5, p=p)
```

power.signrank	<i>Compute power of the signed-rank test</i>
----------------	--

Description

Use the Noether (1987) formula to compute the power of the signed-rank test, which is based on asymptotic normal approximation.

Usage

```
power.signrank(n, alpha, p1, p2)
```

Arguments

n	sample size (scalar)
alpha	p-value threshold (scalar)
p1	$\Pr(X > 0)$, as in Noether (JASA 1987)
p2	$\Pr(X + X' > 0)$, as in Noether (JASA 1987)

Details

In most applications, the null effect size will be designated by $p1 = p2 = 0.5$

Value

Vector of power estimates for two-sided tests

References

Noether, Gottfried E (1987) Sample size determination for some common nonparametric tests. Journal of the American Statistical Association, 82:645-647.

Examples

```
p1 = rep(c(0.8, 0.5), c(100, 900));
p2 = rep(c(0.8, 0.5), c(100, 900));
res = power.signrank(n = 50, alpha = 0.05, p1 = p1, p2 = p2)
```

power.signtest	<i>Compute power of the sign test</i>
----------------	---------------------------------------

Description

Use the Noether (1987) formula to compute the power of the sign test, which is based on asymptotic normal approximation.

Usage

```
power.signtest(n, alpha, p)
```

Arguments

n	sample size (scalar)
alpha	p-value threshold (scalar)
p	$\Pr(X>0)$, as in Noether (JASA 1987)

Details

In most applications, the null effect size will be designated by $p = 0.5$

Value

Vector of power estimates for two-sided tests

References

Noether, Gottfried E (1987) Sample size determination for some common nonparametric tests. Journal of the American Statistical Association, 82:645-647.

Examples

```
p = rep(c(0.8,0.5),c(100,900));  
res = power.signtest(n = 50, alpha = 0.05, p = p)
```

power.tcorr	<i>Compute power of the t-test for non-zero correlation</i>
-------------	---

Description

Compute power of the t-test for non-zero correlation

Usage

```
power.tcorr(n, alpha, rho)
```

Arguments

n	sample size (scalar)
alpha	p-value threshold (scalar)
rho	population correlation coefficient (vector)

Details

For many applications, the null.effect is $\rho = 0$

Value

Vector of power estimates for two-sided tests

Examples

```
rho = rep(c(0.3,0),c(100,900));  
res = power.tcorr(n = 50, alpha = 0.05, rho = rho)
```

Index

`alpha.power.fdr`, [2](#)
`average.power.coxph`, [3](#)
`average.power.fisher`, [4](#)
`average.power.hart`, [5](#)
`average.power.li`, [6](#)
`average.power.oneway`, [7](#)
`average.power.ranksum`, [8](#)
`average.power.signrank`, [8](#)
`average.power.signtest`, [9](#)
`average.power.t.test`, [10](#)
`average.power.tcorr`, [11](#)
`average.power.twoprop`, [11](#)

`fdr.avepow`, [12](#)
`fdr.power.alpha`, [13](#)
`find.sample.size`, [14](#)

`n.fdr.coxph`, [16](#)
`n.fdr.fisher`, [17](#)
`n.fdr.negbin`, [18](#)
`n.fdr.oneway`, [19](#)
`n.fdr.poisson`, [20](#)
`n.fdr.ranksum`, [21](#)
`n.fdr.signrank`, [22](#)
`n.fdr.signtest`, [23](#)
`n.fdr.tcorr`, [24](#)
`n.fdr.ttest`, [25](#)
`n.fdr.twoprop`, [27](#)

`power.cox`, [4](#), [28](#)
`power.fisher`, [5](#), [29](#)
`power.hart`, [6](#), [29](#)
`power.li`, [7](#), [30](#)
`power.oneway`, [7](#), [31](#)
`power.ranksum`, [8](#), [32](#)
`power.signrank`, [9](#), [33](#)
`power.signtest`, [9](#), [34](#)
`power.tcorr`, [11](#), [35](#)