



Moderated estimation of fold change and dispersion for RNA-seq data with DESeq2

Citation

Love, Michael I, Wolfgang Huber, and Simon Anders. 2014. "Moderated estimation of fold change and dispersion for RNA-seq data with DESeq2." *Genome Biology* 15 (12): 550.
doi:10.1186/s13059-014-0550-8. <http://dx.doi.org/10.1186/s13059-014-0550-8>.

Published version

<https://doi.org/10.1186/s13059-014-0550-8>

Link

<http://nrs.harvard.edu/urn-3:HUL.InstRepos:14065371>

Terms of use

This article was downloaded from Harvard University's DASH repository, and is made available under the terms and conditions applicable to Other Posted Material (LAA), as set forth at

<https://harvardwiki.atlassian.net/wiki/external/NGY5NDE4ZjgzNTc5NDQzMGIzZWZhMGFIOWI2M2EwYTg>

Accessibility

<https://accessibility.huit.harvard.edu/digital-accessibility-policy>

Share Your Story

The Harvard community has made this article openly available.
Please share how this access benefits you. [Submit a story](#)

METHOD

Open Access

Moderated estimation of fold change and dispersion for RNA-seq data with DESeq2

Michael I Love^{1,2,3}, Wolfgang Huber² and Simon Anders^{2*}

Abstract

In comparative high-throughput sequencing assays, a fundamental task is the analysis of count data, such as read counts per gene in RNA-seq, for evidence of systematic changes across experimental conditions. Small replicate numbers, discreteness, large dynamic range and the presence of outliers require a suitable statistical approach. We present *DESeq2*, a method for differential analysis of count data, using shrinkage estimation for dispersions and fold changes to improve stability and interpretability of estimates. This enables a more quantitative analysis focused on the strength rather than the mere presence of differential expression. The *DESeq2* package is available at <http://www.bioconductor.org/packages/release/bioc/html/DESeq2.html>.

Background

The rapid adoption of high-throughput sequencing (HTS) technologies for genomic studies has resulted in a need for statistical methods to assess quantitative differences between experiments. An important task here is the analysis of RNA sequencing (RNA-seq) data with the aim of finding genes that are differentially expressed across groups of samples. This task is general: methods for it are typically also applicable for other comparative HTS assays, including chromatin immunoprecipitation sequencing, chromosome conformation capture, or counting observed taxa in metagenomic studies.

Besides the need to account for the specifics of count data, such as non-normality and a dependence of the variance on the mean, a core challenge is the small number of samples in typical HTS experiments – often as few as two or three replicates per condition. Inferential methods that treat each gene separately suffer here from lack of power, due to the high uncertainty of within-group variance estimates. In high-throughput assays, this limitation can be overcome by pooling information across genes, specifically, by exploiting assumptions about the similarity of the variances of different genes measured in the same experiment [1].

Many methods for differential expression analysis of RNA-seq data perform such information sharing across genes for variance (or, equivalently, dispersion) estimation. *edgeR* [2,3] moderates the dispersion estimate for each gene toward a common estimate across all genes, or toward a local estimate from genes with similar expression strength, using a weighted conditional likelihood. Our *DESeq* method [4] detects and corrects dispersion estimates that are too low through modeling of the dependence of the dispersion on the average expression strength over all samples. *BSeq* [5] models the dispersion on the mean, with the mean absolute deviation of dispersion estimates used to reduce the influence of outliers. *DSS* [6] uses a Bayesian approach to provide an estimate for the dispersion for individual genes that accounts for the heterogeneity of dispersion values for different genes. *baySeq* [7] and *ShrinkBayes* [8] estimate priors for a Bayesian model over all genes, and then provide posterior probabilities or false discovery rates (FDRs) for differential expression.

The most common approach in the comparative analysis of transcriptomics data is to test the null hypothesis that the logarithmic fold change (LFC) between treatment and control for a gene's expression is exactly zero, i.e., that the gene is not at all affected by the treatment. Often the goal of differential analysis is to produce a list of genes passing multiple-test adjustment, ranked by *P* value. However, small changes, even if statistically highly significant, might not be the most interesting candidates for

*Correspondence: sanders@fs.tum.de

²Genome Biology Unit, European Molecular Biology Laboratory, Meyerhofstrasse 1, 69117 Heidelberg, Germany
Full list of author information is available at the end of the article

further investigation. Ranking by fold change, on the other hand, is complicated by the noisiness of LFC estimates for genes with low counts. Furthermore, the number of genes called significantly differentially expressed depends as much on the sample size and other aspects of experimental design as it does on the biology of the experiment – and well-powered experiments often generate an overwhelmingly long list of hits [9]. We, therefore, developed a statistical framework to facilitate gene ranking and visualization based on stable estimation of effect sizes (LFCs), as well as testing of differential expression with respect to user-defined thresholds of biological significance.

Here we present *DESeq2*, a successor to our *DESeq* method [4]. *DESeq2* integrates methodological advances with several novel features to facilitate a more quantitative analysis of comparative RNA-seq data using shrinkage estimators for dispersion and fold change. We demonstrate the advantages of *DESeq2*'s new features by describing a number of applications possible with shrunken fold changes and their estimates of standard error, including improved gene ranking and visualization, hypothesis tests above and below a threshold, and the regularized logarithm transformation for quality assessment and clustering of overdispersed count data. We furthermore compare *DESeq2*'s statistical power with existing tools, revealing that our methodology has high sensitivity and precision, while controlling the false positive rate. *DESeq2* is available [10] as an R/Bioconductor package [11].

Results and discussion

Model and normalization

The starting point of a *DESeq2* analysis is a count matrix K with one row for each gene i and one column for each sample j . The matrix entries K_{ij} indicate the number of sequencing reads that have been unambiguously mapped to a gene in a sample. Note that although we refer in this paper to counts of reads in genes, the methods presented here can be applied as well to other kinds of HTS count data. For each gene, we fit a generalized linear model (GLM) [12] as follows.

We model read counts K_{ij} as following a negative binomial distribution (sometimes also called a gamma-Poisson distribution) with mean μ_{ij} and dispersion α_i . The mean is taken as a quantity q_{ij} , proportional to the concentration of cDNA fragments from the gene in the sample, scaled by a normalization factor s_{ij} , i.e., $\mu_{ij} = s_{ij}q_{ij}$. For many applications, the same constant s_j can be used for all genes in a sample, which then accounts for differences in sequencing depth between samples. To estimate these *size factors*, the *DESeq2* package offers the median-of-ratios method already used in *DESeq* [4]. However, it can be advantageous to calculate gene-specific normalization factors s_{ij} to account for further sources of technical biases such as differing dependence on GC content, gene length or the

like, using published methods [13,14], and these can be supplied instead.

We use GLMs with a logarithmic link, $\log_2 q_{ij} = \sum_r x_{jr} \beta_{ir}$, with design matrix elements x_{jr} and coefficients β_{ir} . In the simplest case of a comparison between two groups, such as treated and control samples, the design matrix elements indicate whether a sample j is treated or not, and the GLM fit returns coefficients indicating the overall expression strength of the gene and the $\log_2$ fold change between treatment and control. The use of linear models, however, provides the flexibility to also analyze more complex designs, as is often useful in genomic studies [15].

Empirical Bayes shrinkage for dispersion estimation

Within-group variability, i.e., the variability between replicates, is modeled by the dispersion parameter α_i , which describes the variance of counts via $\text{Var } K_{ij} = \mu_{ij} + \alpha_i \mu_{ij}^2$. Accurate estimation of the dispersion parameter α_i is critical for the statistical inference of differential expression. For studies with large sample sizes this is usually not a problem. For controlled experiments, however, sample sizes tend to be smaller (experimental designs with as little as two or three replicates are common and reasonable), resulting in highly variable dispersion estimates for each gene. If used directly, these noisy estimates would compromise the accuracy of differential expression testing.

One sensible solution is to share information across genes. In *DESeq2*, we assume that genes of similar average expression strength have similar dispersion. We here explain the concepts of our approach using as examples a dataset by Bottomly *et al.* [16] with RNA-seq data for mice of two different strains and a dataset by Pickrell *et al.* [17] with RNA-seq data for human lymphoblastoid cell lines. For the mathematical details, see Methods.

We first treat each gene separately and estimate gene-wise dispersion estimates (using maximum likelihood), which rely only on the data of each individual gene (black dots in Figure 1). Next, we determine the location parameter of the distribution of these estimates; to allow for dependence on average expression strength, we fit a smooth curve, as shown by the red line in Figure 1. This provides an accurate estimate for the expected dispersion value for genes of a given expression strength but does not represent deviations of individual genes from this overall trend. We then shrink the gene-wise dispersion estimates toward the values predicted by the curve to obtain final dispersion values (blue arrow heads). We use an empirical Bayes approach (Methods), which lets the strength of shrinkage depend (i) on an estimate of how close true dispersion values tend to be to the fit and (ii) on the degrees of freedom: as the sample size increases, the shrinkage decreases in strength, and eventually becomes negligible. Our approach therefore accounts for gene-specific

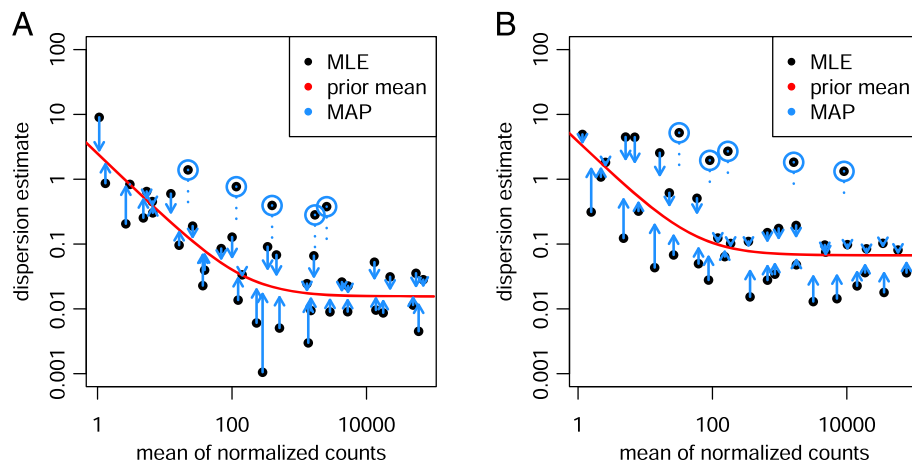


Figure 1 Shrinkage estimation of dispersion. Plot of dispersion estimates over the average expression strength **(A)** for the Bottomly *et al.* [16] dataset with six samples across two groups and **(B)** for five samples from the Pickrell *et al.* [17] dataset, fitting only an intercept term. First, gene-wise MLEs are obtained using only the respective gene's data (black dots). Then, a curve (red) is fit to the MLEs to capture the overall trend of dispersion-mean dependence. This fit is used as a prior mean for a second estimation round, which results in the final MAP estimates of dispersion (arrow heads). This can be understood as a shrinkage (along the blue arrows) of the noisy gene-wise estimates toward the consensus represented by the red line. The black points circled in blue are detected as dispersion outliers and not shrunk toward the prior (shrinkage would follow the dotted line). For clarity, only a subset of genes is shown, which is enriched for dispersion outliers. Additional file 1: Figure S1 displays the same data but with dispersions of all genes shown. MAP, maximum *a posteriori*; MLE, maximum-likelihood estimate.

variation to the extent that the data provide this information, while the fitted curve aids estimation and testing in less information-rich settings.

Our approach is similar to the one used by *DSS* [6], in that both methods sequentially estimate a prior distribution for the true dispersion values around the fit, and then provide the maximum *a posteriori* (MAP) as the final estimate. It differs from the previous implementation of *DESeq*, which used the maximum of the fitted curve and the gene-wise dispersion estimate as the final estimate and tended to overestimate the dispersions (Additional file 1: Figure S2). The approach of *DESeq2* differs from that of *edgeR* [3], as *DESeq2* estimates the width of the prior distribution from the data and therefore automatically controls the amount of shrinkage based on the observed properties of the data. In contrast, the default steps in *edgeR* require a user-adjustable parameter, the *prior degrees of freedom*, which weighs the contribution of the individual gene estimate and *edgeR*'s dispersion fit.

Note that in Figure 1 a number of genes with gene-wise dispersion estimates below the curve have their final estimates raised substantially. The shrinkage procedure thereby helps avoid potential false positives, which can result from underestimates of dispersion. If, on the other hand, an individual gene's dispersion is far above the distribution of the gene-wise dispersion estimates of other genes, then the shrinkage would lead to a greatly reduced final estimate of dispersion. We reasoned that in many cases, the reason for extraordinarily high dispersion of a

gene is that it does not obey our modeling assumptions; some genes may show much higher variability than others for biological or technical reasons, even though they have the same average expression levels. In these cases, inference based on the shrunk dispersion estimates could lead to undesirable false positive calls. *DESeq2* handles these cases by using the gene-wise estimate instead of the shrunk estimate when the former is more than 2 residual standard deviations above the curve.

Empirical Bayes shrinkage for fold-change estimation

A common difficulty in the analysis of HTS data is the strong variance of LFC estimates for genes with low read count. We demonstrate this issue using the dataset by Bottomly *et al.* [16]. As visualized in Figure 2A, weakly expressed genes seem to show much stronger differences between the compared mouse strains than strongly expressed genes. This phenomenon, seen in most HTS datasets, is a direct consequence of dealing with *count* data, in which ratios are inherently noisier when counts are low. This heteroskedasticity (variance of LFCs depending on mean count) complicates downstream analysis and data interpretation, as it makes effect sizes difficult to compare across the dynamic range of the data.

DESeq2 overcomes this issue by shrinking LFC estimates toward zero in a manner such that shrinkage is stronger when the available information for a gene is low, which may be because counts are low, dispersion is high or there are few degrees of freedom. We again employ an empirical Bayes procedure: we first perform

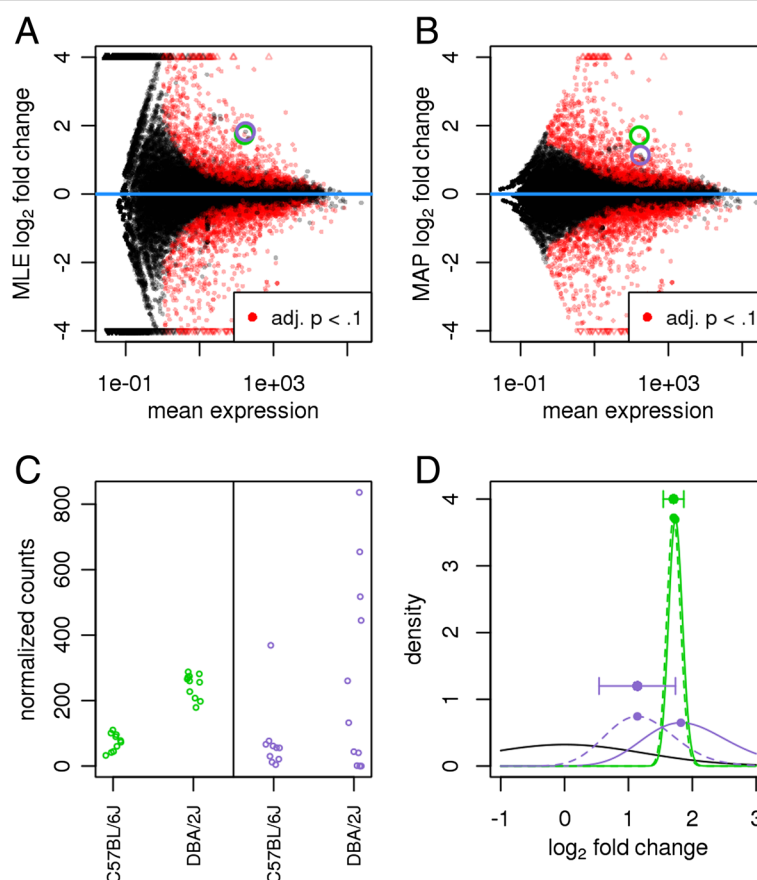


Figure 2 Effect of shrinkage on logarithmic fold change estimates. Plots of the **(A)** MLE (i.e., no shrinkage) and **(B)** MAP estimate (i.e., with shrinkage) for the LFCs attributable to mouse strain, over the average expression strength for a ten vs eleven sample comparison of the Bottomly *et al.* [16] dataset. Small triangles at the top and bottom of the plots indicate points that would fall outside of the plotting window. Two genes with similar mean count and MLE logarithmic fold change are highlighted with green and purple circles. **(C)** The counts (normalized by size factors s_j) for these genes reveal low dispersion for the gene in green and high dispersion for the gene in purple. **(D)** Density plots of the likelihoods (solid lines, scaled to integrate to 1) and the posteriors (dashed lines) for the green and purple genes and of the prior (solid black line): due to the higher dispersion of the purple gene, its likelihood is wider and less peaked (indicating less information), and the prior has more influence on its posterior than for the green gene. The stronger curvature of the green posterior at its maximum translates to a smaller reported standard error for the MAP LFC estimate (horizontal error bar). adj., adjusted; LFC, logarithmic fold change; MAP, maximum *a posteriori*; MLE, maximum-likelihood estimate.

ordinary GLM fits to obtain maximum-likelihood estimates (MLEs) for the LFCs and then fit a zero-centered normal distribution to the observed distribution of MLEs over all genes. This distribution is used as a prior on LFCs in a second round of GLM fits, and the MAP estimates are kept as final estimates of LFC. Furthermore, a standard error for each estimate is reported, which is derived from the posterior's curvature at its maximum (see Methods for details). These shrunk LFCs and their standard errors are used in the Wald tests for differential expression described in the next section.

The resulting MAP LFCs are biased toward zero in a manner that removes the problem of exaggerated LFCs for low counts. As Figure 2B shows, the strongest LFCs are no longer exhibited by genes with weakest expression. Rather, the estimates are more evenly spread around zero, and

for very weakly expressed genes (with less than one read per sample on average), LFCs hardly deviate from zero, reflecting that accurate LFC estimates are not possible here.

The strength of shrinkage does not depend simply on the mean count, but rather on the amount of information available for the fold change estimation (as indicated by the observed Fisher information; see Methods). Two genes with equal expression strength but different dispersions will experience a different amount of shrinkage (Figure 2C,D). The shrinkage of LFC estimates can be described as a bias-variance trade-off [18]: for genes with little information for LFC estimation, a reduction of the strong variance is bought at the cost of accepting a bias toward zero, and this can result in an overall reduction in mean squared error, e.g., when comparing to LFC

estimates from a new dataset. Genes with high information for LFC estimation will have, in our approach, LFCs with both low bias and low variance. Furthermore, as the degrees of freedom increase, and the experiment provides more information for LFC estimation, the shrunken estimates will converge to the unshrunk estimates. We note that other Bayesian efforts toward moderating fold changes for RNA-seq include hierarchical models [8,19] and the *GFOLD* (or generalized fold change) tool [20], which uses a posterior distribution of LFCs.

The shrunken MAP LFCs offer a more reproducible quantification of transcriptional differences than standard MLE LFCs. To demonstrate this, we split the Bottomly *et al.* samples equally into two groups, I and II, such that each group contained a balanced split of the strains, simulating a scenario where an experiment (samples in group I) is performed, analyzed and reported, and then independently replicated (samples in group II). Within each group, we estimated LFCs between the strains and compared between groups I and II, using the MLE LFCs (Figure 3A) and using the MAP LFCs (Figure 3B). Because the shrinkage moves large LFCs that are not well supported by the data toward zero, the agreement between the two independent sample groups increases considerably. Therefore, shrunken fold-change estimates offer a more reliable basis for quantitative conclusions than normal MLEs.

This makes shrunken LFCs also suitable for ranking genes, e.g., to prioritize them for follow-up experiments. For example, if we sort the genes in the two sample groups of Figure 3 by unshrunk LFC estimates, and consider the 100 genes with the strongest up- or down-regulation in group I, we find only 21 of these again among the top 100 up- or down-regulated genes in group II. However, if

we rank the genes by shrunken LFC estimates, the overlap improves to 81 of 100 genes (Additional file 1: Figure S3).

A simpler often used method is to add a fixed number (pseudocount) to all counts before forming ratios. However, this requires the choice of a tuning parameter and only reacts to one of the sources of uncertainty, low counts, but not to gene-specific dispersion differences or sample size. We demonstrate this in the *Benchmarks* section below.

Hypothesis tests for differential expression

After GLMs are fit for each gene, one may test whether each model coefficient differs significantly from zero. *DESeq2* reports the standard error for each shrunken LFC estimate, obtained from the curvature of the coefficient's posterior (dashed lines in Figure 2D) at its maximum. For significance testing, *DESeq2* uses a Wald test: the shrunken estimate of LFC is divided by its standard error, resulting in a *z*-statistic, which is compared to a standard normal distribution. (See Methods for details.) The Wald test allows testing of individual coefficients, or contrasts of coefficients, without the need to fit a reduced model as with the likelihood ratio test, though the likelihood ratio test is also available as an option in *DESeq2*. The Wald test *P* values from the subset of genes that pass an independent filtering step, described in the next section, are adjusted for multiple testing using the procedure of Benjamini and Hochberg [21].

Automatic independent filtering

Due to the large number of tests performed in the analysis of RNA-seq and other genome-wide experiments, the multiple testing problem needs to be addressed. A popular objective is control or estimation of the FDR. Multiple

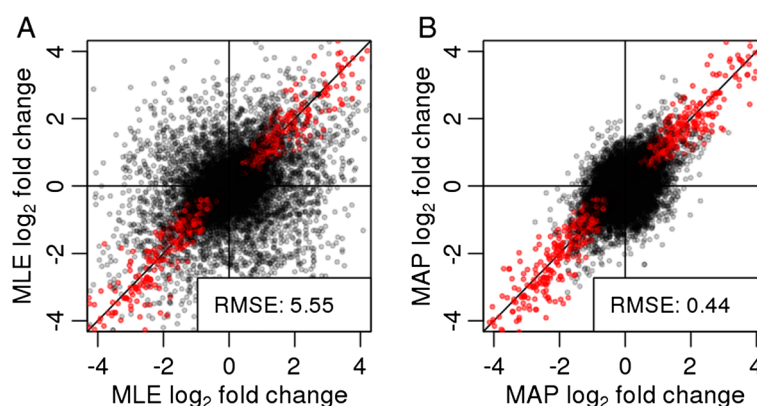


Figure 3 Stability of logarithmic fold changes. *DESeq2* is run on equally split halves of the data of Bottomly *et al.* [16], and the LFCs from the halves are plotted against each other. **(A)** MLEs, i.e., without LFC shrinkage. **(B)** MAP estimates, i.e., with shrinkage. Points in the top left and bottom right quadrants indicate genes with a change of sign of LFC. Red points indicate genes with adjusted *P* value < 0.1. The legend displays the root-mean-square error of the estimates in group I compared to those in group II. LFC, logarithmic fold change; MAP, maximum *a posteriori*; MLE, maximum-likelihood estimate; RMSE, root-mean-square error.

testing adjustment tends to be associated with a loss of power, in the sense that the FDR for a set of genes is often higher than the individual P values of these genes. However, the loss can be reduced if genes that have little or no chance of being detected as differentially expressed are omitted from the testing, provided that the criterion for omission is independent of the test statistic under the null hypothesis [22] (see Methods). *DESeq2* uses the average expression strength of each gene, across all samples, as its filter criterion, and it omits all genes with mean normalized counts below a filtering threshold from multiple testing adjustment. *DESeq2* by default will choose a threshold that maximizes the number of genes found at a user-specified target FDR. In Figures 2A,B and 3, genes found in this way to be significant at an estimated FDR of 10% are depicted in red. Depending on the distribution of the mean normalized counts, the resulting increase in power can be substantial, sometimes making the difference in whether or not any differentially expressed genes are detected.

Hypothesis tests with thresholds on effect size

Specifying minimum effect size

Most approaches to testing for differential expression, including the default approach of *DESeq2*, test against the null hypothesis of *zero* LFC. However, if any biological processes are genuinely affected by the difference in experimental treatment, this null hypothesis implies that the gene under consideration is *perfectly* decoupled from these processes. Due to the high interconnectedness of cells' regulatory networks, this hypothesis is, in fact, implausible, and arguably wrong for many if not most genes. Consequently, with sufficient sample size, even genes with a very small but non-zero LFC will eventually be detected as differentially expressed. A change should therefore be of sufficient magnitude to be considered *biologically significant*. For small-scale experiments, statistical significance is often a much stricter requirement than biological significance, thereby relieving the researcher from the need to decide on a threshold for biological significance.

For well-powered experiments, however, a statistical test against the conventional null hypothesis of zero LFC may report genes with statistically significant changes that are so weak in effect strength that they could be considered irrelevant or distracting. A common procedure is to disregard genes whose estimated LFC β_{ir} is below some threshold, $|\beta_{ir}| \leq \theta$. However, this approach loses the benefit of an easily interpretable FDR, as the reported P value and adjusted P value still correspond to the test of *zero* LFC. It is therefore desirable to include the threshold in the statistical testing procedure directly, i.e., not to filter post hoc on a reported fold-change *estimate*, but rather to evaluate statistically directly whether there

is sufficient evidence that the LFC is above the chosen threshold.

DESeq2 offers tests for composite null hypotheses of the form $|\beta_{ir}| \leq \theta$, where β_{ir} is the shrunken LFC from the estimation procedure described above. (See Methods for details.) Figure 4A demonstrates how such a thresholded test gives rise to a curved decision boundary: to reach significance, the estimated LFC has to exceed the specified threshold by an amount that depends on the available information. We note that related approaches to generate gene lists that satisfy both statistical and biological significance criteria have been previously discussed for microarray data [23] and recently for sequencing data [19].

Specifying maximum effect size

Sometimes, a researcher is interested in finding genes that are not, or only very weakly, affected by the treatment or experimental condition. This amounts to a setting similar to the one just discussed, but the roles of the null and alternative hypotheses are swapped. We are here asking for evidence of the effect being weak, not for evidence of the effect being zero, because the latter question is rarely tractable. The meaning of *weak* needs to be quantified for the biological question at hand by choosing a suitable threshold θ for the LFC. For such analyses, *DESeq2* offers a test of the composite null hypothesis $|\beta_{ir}| \geq \theta$, which will report genes as significant for which there is evidence that their LFC is weaker than θ . Figure 4B shows the outcome of such a test. For genes with very low read count, even an estimate of zero LFC is not significant, as the large uncertainty of the estimate does not allow us to exclude that the gene may in truth be more than weakly affected by the experimental condition. Note the lack of LFC shrinkage: to find genes with weak differential expression, *DESeq2* requires that the LFC shrinkage has been disabled. This is because the zero-centered prior used for LFC shrinkage embodies a *prior* belief that LFCs tend to be small, and hence is inappropriate here.

Detection of count outliers

Parametric methods for detecting differential expression can have gene-wise estimates of LFC overly influenced by individual outliers that do not fit the distributional assumptions of the model [24]. An example of such an outlier would be a gene with single-digit counts for all samples, except one sample with a count in the thousands. As the aim of differential expression analysis is typically to find *consistently* up- or down-regulated genes, it is useful to consider diagnostics for detecting individual observations that overly influence the LFC estimate and P value for a gene. A standard outlier diagnostic is Cook's distance [25], which is defined within each gene for each sample as the scaled distance that the coefficient vector,

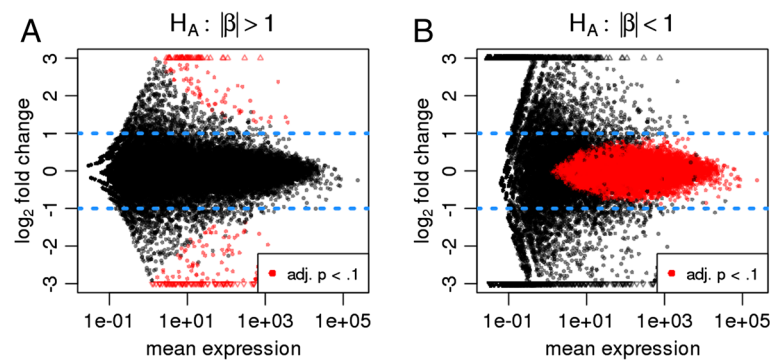


Figure 4 Hypothesis testing involving non-zero thresholds. Shown are plots of the estimated fold change over average expression strength ("minus over average", or MA-plots) for a ten vs eleven comparison using the Bottomly *et al.* [16] dataset, with highlighted points indicating low adjusted P values. The alternate hypotheses are that logarithmic (base 2) fold changes are **(A)** greater than 1 in absolute value or **(B)** less than 1 in absolute value. adj., adjusted.

$\tilde{\beta}_i$, of a linear model or GLM would move if the sample were removed and the model refit.

DESeq2 flags, for each gene, those samples that have a Cook's distance greater than the 0.99 quantile of the $F(p, m - p)$ distribution, where p is the number of model parameters including the intercept, and m is the number of samples. The use of the F distribution is motivated by the heuristic reasoning that removing a single sample should not move the vector $\tilde{\beta}_i$ outside of a 99% confidence region around $\tilde{\beta}_i$ fit using all the samples [25]. However, if there are two or fewer replicates for a condition, these samples do not contribute to outlier detection, as there are insufficient replicates to determine outlier status.

How should one deal with flagged outliers? In an experiment with many replicates, discarding the outlier and proceeding with the remaining data might make best use of the available data. In a small experiment with few samples, however, the presence of an outlier can impair inference regarding the affected gene, and merely ignoring the outlier may even be considered data cherry-picking – and therefore, it is more prudent to exclude the whole gene from downstream analysis.

Hence, *DESeq2* offers two possible responses to flagged outliers. By default, outliers in conditions with six or fewer replicates cause the whole gene to be flagged and removed from subsequent analysis, including P value adjustment for multiple testing. For conditions that contain seven or more replicates, *DESeq2* replaces the outlier counts with an imputed value, namely the trimmed mean over all samples, scaled by the size factor, and then re-estimates the dispersion, LFCs and P values for these genes. As the outlier is replaced with the value predicted by the null hypothesis of no differential expression, this is a more conservative choice than simply omitting the outlier. When there are many degrees of freedom, the second approach avoids discarding genes that might contain true differential expression.

Additional file 1: Figure S4 displays the outlier replacement procedure for a single gene in a seven by seven comparison of the Bottomly *et al.* [16] dataset. While the original fitted means are heavily influenced by a single sample with a large count, the corrected LFCs provide a better fit to the majority of the samples.

Regularized logarithm transformation

For certain analyses, it is useful to transform data to render them homoskedastic. As an example, consider the task of assessing sample similarities in an unsupervised manner using a clustering or ordination algorithm. For RNA-seq data, the problem of heteroskedasticity arises: if the data are given to such an algorithm on the original count scale, the result will be dominated by highly expressed, highly variable genes; if logarithm-transformed data are used, undue weight will be given to weakly expressed genes, which show exaggerated LFCs, as discussed above. Therefore, we use the shrinkage approach of *DESeq2* to implement a *regularized logarithm* transformation (rlog), which behaves similarly to a $\log_2$ transformation for genes with high counts, while shrinking together the values for different samples for genes with low counts. It therefore avoids a commonly observed property of the standard logarithm transformation, the spreading apart of data for genes with low counts, where random noise is likely to dominate any biologically meaningful signal. When we consider the variance of each gene, computed across samples, these variances are stabilized – i.e., approximately the same, or homoskedastic – after the rlog transformation, while they would otherwise strongly depend on the mean counts. It thus facilitates multivariate visualization and ordinations such as clustering or principal component analysis that tend to work best when the variables have similar dynamic range. Note that while the rlog transformation builds upon our LFC shrinkage approach, it is distinct from and not part of the statistical inference

procedure for differential expression analysis described above, which employs the raw counts, not transformed data.

The rlog transformation is calculated by fitting for each gene a GLM with a baseline expression (i.e., intercept only) and, computing for each sample, shrunk LFCs with respect to the baseline, using the same empirical Bayes procedure as before (Methods). Here, however, the sample covariate information (e.g. treatment or control) is not used, so that all samples are treated equally. The rlog transformation accounts for variation in sequencing depth across samples as it represents the logarithm of q_{ij} after accounting for the size factors s_{ij} . This is in contrast to the variance-stabilizing transformation (VST) for overdispersed counts introduced in *DESeq* [4]: while the VST is also effective at stabilizing variance, it does not directly take into account differences in size factors; and in datasets with large variation in sequencing depth (dynamic range of size factors $\gtrsim 4$) we observed undesirable artifacts in the performance of the VST. A disadvantage of the rlog

transformation with respect to the VST is, however, that the ordering of genes within a sample will change if neighboring genes undergo shrinkage of different strength. As with the VST, the value of $\text{rlog}(K_{ij})$ for large counts is approximately equal to $\log_2(K_{ij}/s_j)$. Both the rlog transformation and the VST are provided in the *DESeq2* package.

We demonstrate the use of the rlog transformation on the RNA-seq dataset of Hammer *et al.* [26], wherein RNA was sequenced from the dorsal root ganglion of rats that had undergone spinal nerve ligation and controls, at 2 weeks and at 2 months after the ligation. The count matrix for this dataset was downloaded from the ReCount online resource [27]. This dataset offers more subtle differences between conditions than the Bottomly *et al.* [16] dataset. Figure 5 provides diagnostic plots of the normalized counts under the ordinary logarithm with a pseudocount of 1 and the rlog transformation, showing that the rlog both stabilizes the variance through the range of the mean of counts and helps to find meaningful patterns in the data.

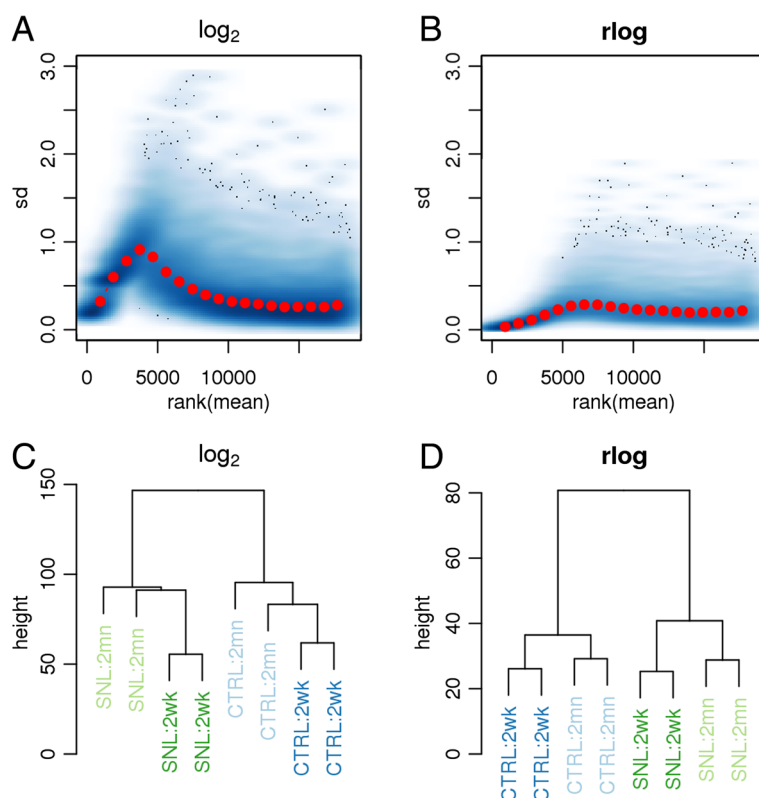


Figure 5 Variance stabilization and clustering after rlog transformation. Two transformations were applied to the counts of the Hammer *et al.* [26] dataset: the logarithm of normalized counts plus a pseudocount, i.e. $f(K_{ij}) = \log_2(K_{ij}/s_j + 1)$, and the rlog. The gene-wise standard deviation of transformed values is variable across the range of the mean of counts using the logarithm (**A**), while relatively stable using the rlog (**B**). A hierarchical clustering on Euclidean distances and complete linkage using the rlog (**D**) transformed data clusters the samples into the groups defined by treatment and time, while using the logarithm-transformed counts (**C**) produces a more ambiguous result. sd, standard deviation.

Gene-level analysis

We here present *DESeq2* for the analysis of per-gene counts, i.e., the total number of reads that can be uniquely assigned to a gene. In contrast, several algorithms [28,29] work with probabilistic assignments of reads to transcripts, where multiple, overlapping transcripts can originate from each gene. It has been noted that the total read count approach can result in false detection of differential expression when in fact only transcript isoform lengths change, and even in a wrong sign of LFCs in extreme cases [28]. However, in our benchmark, discussed in the following section, we found that LFC sign disagreements between total read count and probabilistic-assignment-based methods were rare for genes that were differentially expressed according to either method (Additional file 1: Figure S5). Furthermore, if estimates for average transcript length are available for the conditions, these can be incorporated into the *DESeq2* framework as gene- and sample-specific normalization factors. In addition, the approach used in *DESeq2* can be extended to isoform-specific analysis, either through generalized linear modeling at the exon level with a gene-specific mean as in the *DEXSeq* package [30] or through counting evidence for alternative isoforms in splice graphs [31,32]. In fact, the latest release version of *DEXSeq* now uses *DESeq2* as its inferential engine and so offers shrinkage estimation of dispersion and effect sizes for an exon-level analysis, too.

Comparative benchmarks

To assess how well *DESeq2* performs for standard analyses in comparison to other current methods, we used a combination of simulations and real data. The negative-binomial-based approaches compared were *DESeq (old)* [4], *edgeR* [33], *edgeR* with the robust option [34], *DSS* [6] and *EBSeq* [35]. Other methods compared were the *voom* normalization method followed by linear modeling using the *limma* package [36] and the *SAMseq* permutation method of the *samr* package [24]. For the benchmarks using real data, the *Cuffdiff 2* [28] method of the Cufflinks suite was included. For version numbers of the software used, see Additional file 1: Table S3. For all algorithms returning *P* values, the *P* values from genes with non-zero sum of read counts across samples were adjusted using the Benjamini–Hochberg procedure [21].

Benchmarks through simulation

Sensitivity and precision We simulated datasets of 10,000 genes with negative binomial distributed counts. To simulate data with realistic moments, the mean and dispersions were drawn from the joint distribution of means and gene-wise dispersion estimates from the Pickrell *et al.* data, fitting only an intercept term.

These datasets were of varying total sample size ($m \in \{6, 8, 10, 20\}$), and the samples were split into two equal-sized groups; 80% of the simulated genes had no true differential expression, while for 20% of the genes, true fold changes of 2, 3 and 4 were used to generate counts across the two groups, with the direction of fold change chosen randomly. The simulated differentially expressed genes were chosen uniformly at random among all the genes, throughout the range of mean counts. MA-plots of the true fold changes used in the simulation and the observed fold changes induced by the simulation for one of the simulation settings are shown in Additional file 1: Figure S6.

Algorithms' performance in the simulation benchmark was assessed by their sensitivity and precision. The sensitivity was calculated as the fraction of genes with adjusted *P* value < 0.1 among the genes with true differences between group means. The precision was calculated as the fraction of genes with true differences between group means among those with adjusted *P* value < 0.1 . The sensitivity is plotted over 1 – precision, or the FDR, in Figure 6. *DESeq2*, and also *edgeR*, often had the highest sensitivity of the algorithms that controlled type-I error in the sense that the actual FDR was at or below 0.1, the threshold for adjusted *P* values used for calling differentially expressed genes. *DESeq2* had higher sensitivity compared to the other algorithms, particularly for small fold change (2 or 3), as was also found in benchmarks performed by Zhou *et al.* [34]. For larger sample sizes and larger fold changes the performance of the various algorithms was more consistent.

The overly conservative calling of the old *DESeq* tool can be observed, with reduced sensitivity compared to the other algorithms and an actual FDR less than the nominal value of 0.1. We note that *EBSeq* version 1.4.0 by default removes low-count genes – whose 75% quantile of normalized counts is less than ten – before calling differential expression. The sensitivity of algorithms on the simulated data across a range of the mean of counts are more closely compared in Additional file 1: Figure S9.

Outlier sensitivity We used simulations to compare the sensitivity and specificity of *DESeq2*'s outlier handling approach to that of *edgeR*, which was recently added to the software and published while this manuscript was under review. *edgeR* now includes an optional method to handle outliers by iteratively refitting the GLM after down-weighting potential outlier counts [34]. The simulations, summarized in Additional file 1: Figure S10, indicated that both approaches to outliers nearly recover the performance on an outlier-free dataset, though *edgeR-robust* had slightly higher actual than nominal FDR, as seen in Additional file 1: Figure S11.

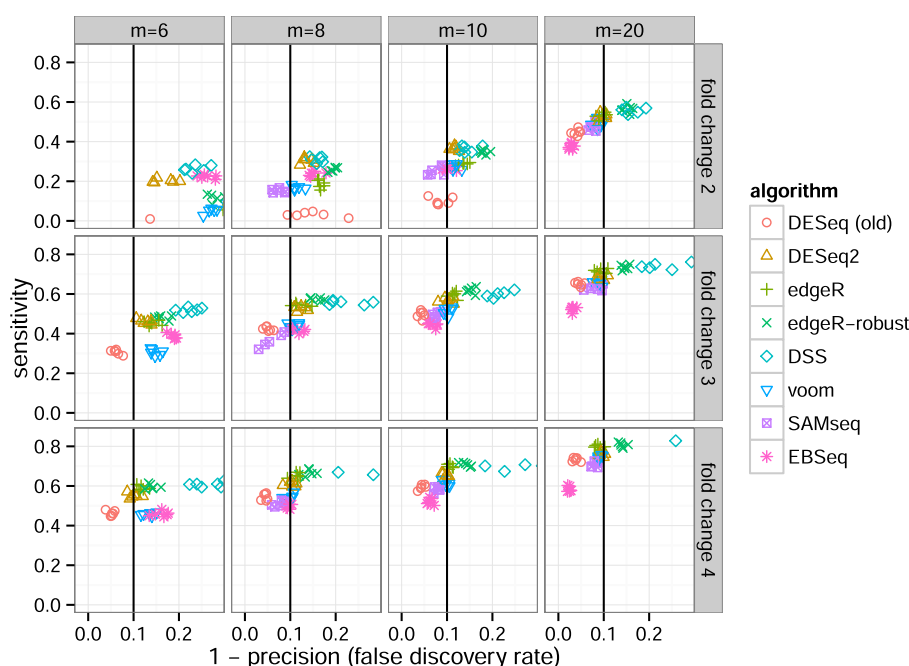


Figure 6 Sensitivity and precision of algorithms across combinations of sample size and effect size. *DESeq2* and *edgeR* often had the highest sensitivity of those algorithms that controlled the FDR, i.e., those algorithms which fall on or to the left of the vertical black line. For a plot of sensitivity against false positive rate, rather than FDR, see Additional file 1: Figure S8, and for the dependence of sensitivity on the mean of counts, see Additional file 1: Figure S9. Note that *EBSeq* filters low-count genes (see main text for details).

Precision of fold change estimates We benchmarked the *DESeq2* approach of using an empirical prior to achieve shrinkage of LFC estimates against two competing approaches: the *GFOLD* method, which can analyze experiments without replication [20] and can also handle experiments with replicates, and the *edgeR* package, which provides a pseudocount-based shrinkage termed *predictive LFCs*. Results are summarized in Additional file 1: Figures S12–S16. *DESeq2* had consistently low root-mean-square error and mean absolute error across a range of sample sizes and models for a distribution of true LFCs. *GFOLD* had similarly low error to *DESeq2* over all genes; however, when focusing on differentially expressed genes, it performed worse for larger sample sizes. *edgeR* with default settings had similarly low error to *DESeq2* when focusing only on the differentially expressed genes, but had higher error over all genes.

Clustering We compared the performance of the rlog transformation against other methods of transformation or distance calculation in the recovery of simulated clusters. The adjusted Rand index [37] was used to compare a hierarchical clustering based on various distances with the true cluster membership. We tested the Euclidean distance for normalized counts, logarithm of normalized counts plus a pseudocount of 1, rlog-transformed counts and VST counts. In addition we compared these Euclidean

distances with the Poisson distance implemented in the *PoiClaChlu* package [38], and a distance implemented internally in the *plotMDS* function of *edgeR* (though not the default distance, which is similar to the logarithm of normalized counts). The results, shown in Additional file 1: Figure S17, revealed that when the size factors were equal for all samples, the Poisson distance and the Euclidean distance of rlog-transformed or VST counts outperformed other methods. However, when the size factors were not equal across samples, the rlog approach generally outperformed the other methods. Finally, we note that the rlog transformation provides normalized data, which can be used for a variety of applications, of which distance calculation is one.

Benchmark for RNA sequencing data

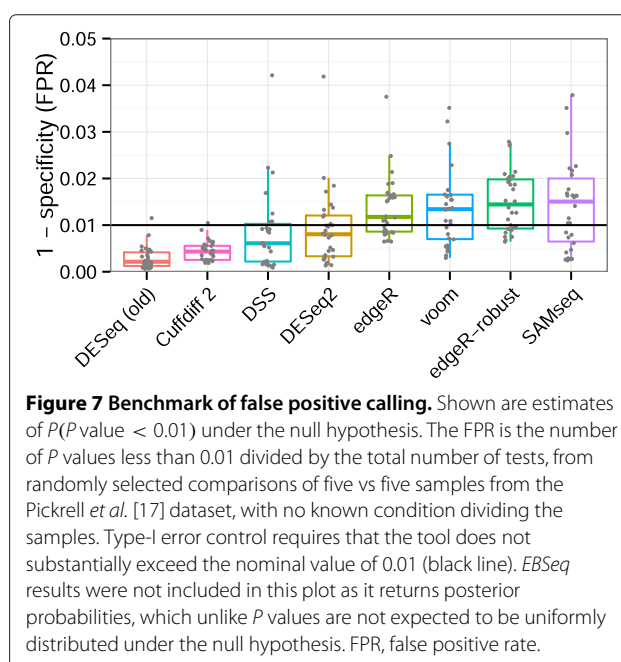
While simulation is useful to verify how well an algorithm behaves with idealized theoretical data, and hence can verify that the algorithm performs as expected under its own assumptions, simulations cannot inform us how well the theory fits reality. With RNA-seq data, there is the complication of not knowing fully or directly the underlying truth; however, we can work around this limitation by using more indirect inference, explained below.

In the following benchmarks, we considered three performance metrics for differential expression calling: the false positive rate (or 1 minus the specificity), sensitivity

and precision. We can obtain meaningful estimates of specificity from looking at datasets where we believe all genes fall under the null hypothesis of no differential expression [39]. Sensitivity and precision are more difficult to estimate, as they require independent knowledge of those genes that are differentially expressed. To circumvent this problem, we used experimental reproducibility on independent samples (though from the same dataset) as a proxy. We used a dataset with large numbers of replicates in both of two groups, where we expect that truly differentially expressed genes exist. We repeatedly split this dataset into an evaluation set and a larger verification set, and compared the calls from the evaluation set with the calls from the verification set, which were taken as truth. It is important to keep in mind that the calls from the verification set are only an approximation of the true differential state, and the approximation error has a systematic and a stochastic component. The stochastic error becomes small once the sample size of the verification set is large enough. For the systematic errors, our benchmark assumes that these affect all algorithms more or less equally and do not markedly change the ranking of the algorithms.

False positive rate To evaluate the false positive rate of the algorithms, we considered mock comparisons from a dataset with many samples and no known condition dividing the samples into distinct groups. We used the RNA-seq data of Pickrell *et al.* [17] for lymphoblastoid cell lines derived from unrelated Nigerian individuals. We chose a set of 26 RNA-seq samples of the same read length (46 base pairs) from male individuals. We randomly drew without replacement ten samples from the set to compare five against five, and this process was repeated 30 times. We estimated the false positive rate associated with a critical value of 0.01 by dividing the number of P values less than 0.01 by the total number of tests; genes with zero sum of read counts across samples were excluded. The results over the 30 replications, summarized in Figure 7, indicated that all algorithms generally controlled the number of false positives. *DESeq (old)* and *Cuffdiff 2* appeared overly conservative in this analysis, not using up their type-I error budget.

Sensitivity To obtain an impression of the sensitivity of the algorithms, we considered the Bottomly *et al.* [16] dataset, which contains ten and eleven replicates of two different, genetically homogeneous mice strains. This allowed for a split of three vs three for the evaluation set and seven vs eight for the verification set, which were balanced across the three experimental batches. Random splits were replicated 30 times. Batch information was not provided to the *DESeq (old)*, *DESeq2*, *DSS*, *edgeR* or *voom* algorithms, which can accommodate complex

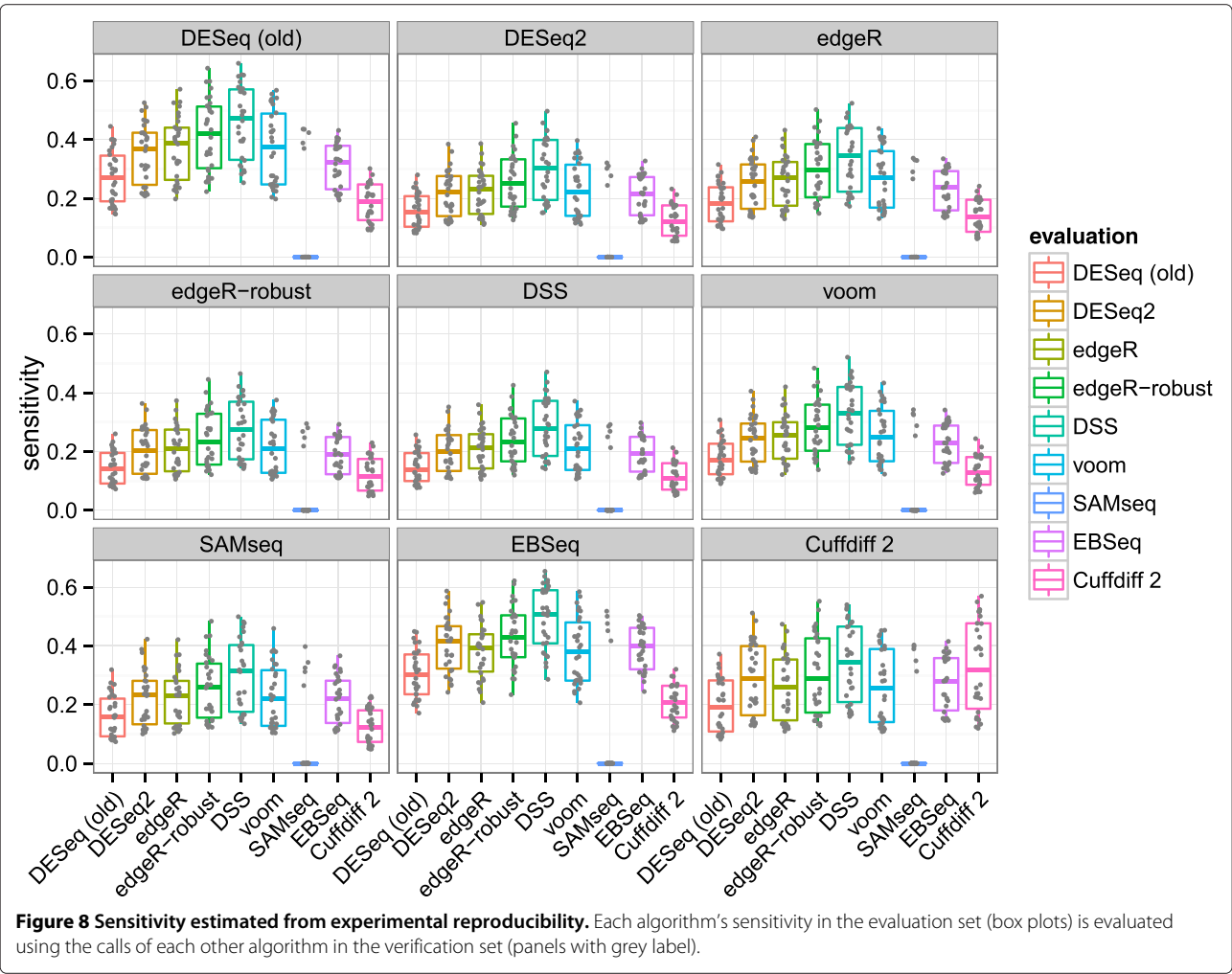


experimental designs, to have comparable calls across all algorithms.

We rotated through each algorithm to determine the calls of the verification set. For a given algorithm's verification set calls, we tested the evaluation set calls of every algorithm. We used this approach rather than a consensus-based method, as we did not want to favor or disfavor any particular algorithm or group of algorithms. Sensitivity was calculated as in the simulation benchmark, now with true differential expression defined by an adjusted P value < 0.1 in the larger verification set, as diagrammed in Additional file 1: Figure S18. Figure 8 displays the estimates of sensitivity for each algorithm pair.

The ranking of algorithms was generally consistent regardless of which algorithm was chosen to determine calls in the verification set. *DESeq2* had comparable sensitivity to *edgeR* and *voom* though less than *DSS*. The median sensitivity estimates were typically between 0.2 and 0.4 for all algorithms. That all algorithms had relatively low median sensitivity can be explained by the small sample size of the evaluation set and the fact that increasing the sample size in the verification set increases power. It was expected that the permutation-based *SAMseq* method would rarely produce adjusted P value < 0.1 in the evaluation set, because the three vs three comparison does not enable enough permutations.

Precision Another important consideration from the perspective of an investigator is the precision, or fraction of true positives in the set of genes which pass the adjusted P value threshold. This can also be reported as



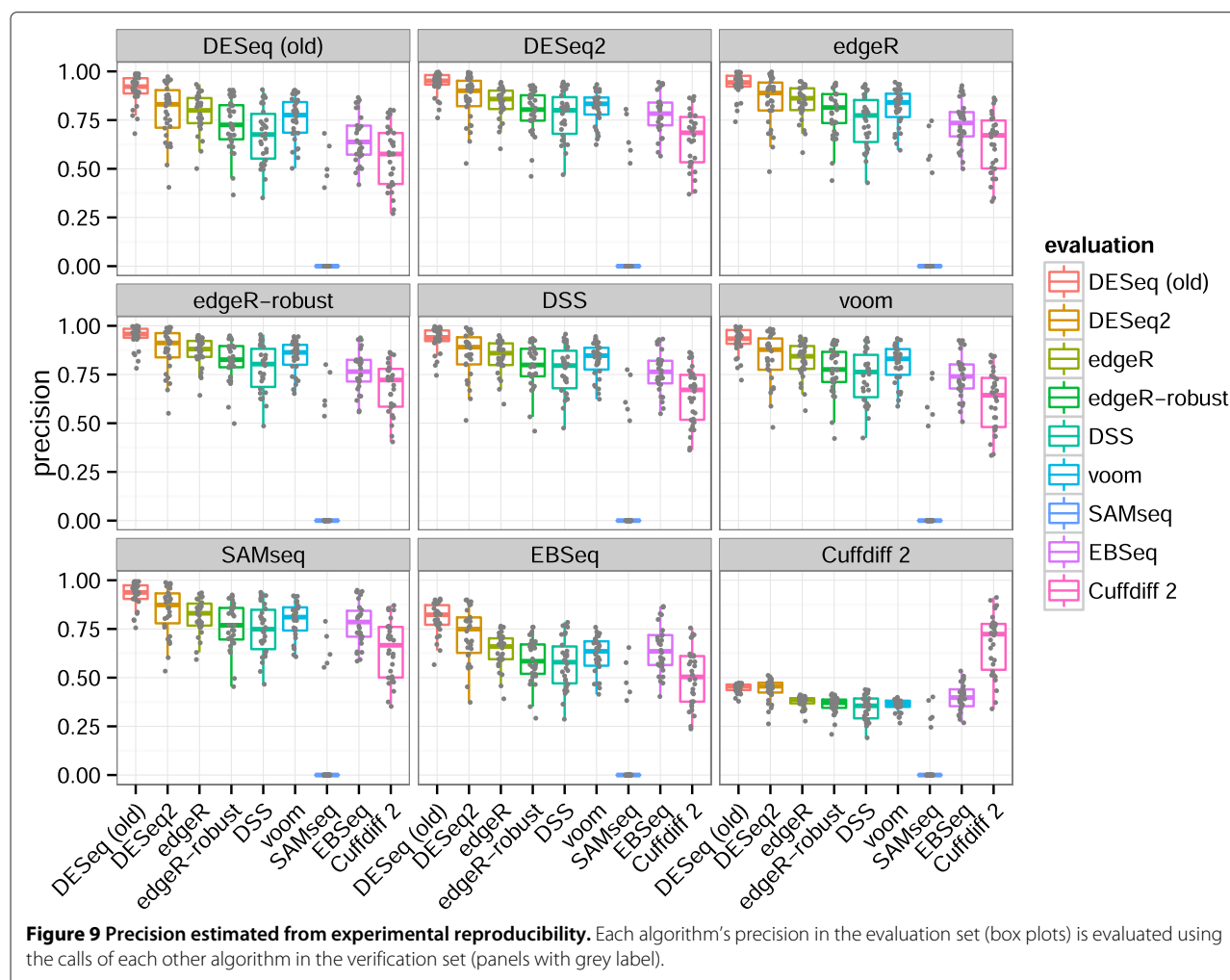
1 – FDR. Again, ‘true’ differential expression was defined by an adjusted P value < 0.1 in the larger verification set. The estimates of precision are displayed in Figure 9, where we can see that *DESeq2* often had the second highest median precision, behind *DESeq (old)*. We can also see that algorithms with higher median sensitivity, e.g., *DSS*, were generally associated here with lower median precision. The rankings differed significantly when *Cuffdiff 2* was used to determine the verification set calls. This is likely due to the additional steps *Cuffdiff 2* performed to deconvolve changes in isoform-level abundance from gene-level abundance, which apparently came at the cost of lower precision when compared against its own verification set calls.

To compare the sensitivity and precision results further, we calculated the precision of algorithms along a grid of nominal adjusted P values (Additional file 1: Figure S19). We then found the nominal adjusted P value for each algorithm, which resulted in a median actual precision of 0.9 (FDR = 0.1). Having thus calibrated each algorithm to

a target FDR, we evaluated the sensitivity of calling, as shown in Additional file 1: Figure S20. As expected, here the algorithms performed more similarly to each other. This analysis revealed that, for a given target precision, *DESeq2* often was among the top algorithms by median sensitivity, though the variability across random replicates was larger than the differences between algorithms.

The absolute number of calls for the evaluation and verification sets can be seen in Additional file 1: Figures S21 and S22, which mostly matched the order seen in the sensitivity plot of Figure 8. Additional file 1: Figures S23 and S24 provide heat maps and clustering based on the Jaccard index of calls for one replicate of the evaluation and verification sets, indicating a large overlap of calls across the different algorithms.

In summary, the benchmarking tests showed that *DESeq2* effectively controlled type-I errors, maintaining a median false positive rate just below the chosen critical value in a mock comparison of groups of samples randomly chosen from a larger pool. For both simulation and



analysis of real data, *DESeq2* often achieved the highest sensitivity of those algorithms that controlled the FDR.

Conclusions

DESeq2 offers a comprehensive and general solution for gene-level analysis of RNA-seq data. Shrinkage estimators substantially improve the stability and reproducibility of analysis results compared to maximum-likelihood-based solutions. Empirical Bayes priors provide automatic control of the amount of shrinkage based on the amount of information for the estimated quantity available in the data. This allows *DESeq2* to offer consistent performance over a large range of data types and makes it applicable for small studies with few replicates as well as for large observational studies. *DESeq2*'s heuristics for outlier detection help to recognize genes for which the modeling assumptions are unsuitable and so avoids type-I errors caused by these. The embedding of these strategies in the framework of GLMs enables the treatment of both simple and complex designs.

A critical advance is the shrinkage estimator for fold changes for differential expression analysis, which offers a sound and statistically well-founded solution to the practically relevant problem of comparing fold change across the wide dynamic range of RNA-seq experiments. This is of value for many downstream analysis tasks, including the ranking of genes for follow-up studies and association of fold changes with other variables of interest. In addition, the rlog transformation, which implements shrinkage of fold changes on a per-sample basis, facilitates visualization of differences, for example in heat maps, and enables the application of a wide range of techniques that require homoskedastic input data, including machine-learning or ordination techniques such as principal component analysis and clustering.

DESeq2 hence offers to practitioners a wide set of features with state-of-the-art inferential power. Its use cases are not limited to RNA-seq data or other transcriptomics assays; rather, many kinds of high-throughput count data can be used. Other areas for which *DESeq* or *DESeq2*

have been used include chromatin immunoprecipitation sequencing assays (e.g., [40]; see also the *DiffBind* package [41,42]), barcode-based assays (e.g., [43]), metagenomics data (e.g., [44]), ribosome profiling [45] and CRISPR/Cas-library assays [46]. Finally, the *DESeq2* package is integrated well in the Bioconductor infrastructure [11] and comes with extensive documentation, including a vignette that demonstrates a complete analysis step by step and discusses advanced use cases.

Materials and methods

A summary of the notation used in the following section is provided in Additional file 1: Table S1.

Model and normalization

The read count K_{ij} for gene i in sample j is described with a GLM of the negative binomial family with a logarithmic link:

$$K_{ij} \sim \text{NB}(\text{mean} = \mu_{ij}, \text{dispersion} = \alpha_i) \quad (1)$$

$$\mu_{ij} = s_{ij}q_{ij}$$

$$\log q_{ij} = \sum_r x_{jr}\beta_{ir}. \quad (2)$$

For notational simplicity, the equations here use the natural logarithm as the link function, though the *DESeq2* software reports estimated model coefficients and their estimated standard errors on the $\log_2$ scale.

By default, the normalization constants s_{ij} are considered constant within a sample, $s_{ij} = s_j$, and are estimated with the median-of-ratios method previously described and used in *DESeq* [4] and *DEXSeq* [30]:

$$s_j = \text{median}_{i: K_i^R \neq 0} \frac{K_{ij}}{K_i^R} \quad \text{with} \quad K_i^R = \left(\prod_{j=1}^m K_{ij} \right)^{1/m}.$$

Alternatively, the user can supply normalization constants s_{ij} calculated using other methods (e.g., using *cqn* [13] or *EDASeq* [14]), which may differ from gene to gene.

Expanded design matrices

For consistency with our software's documentation, in the following text we will use the terminology of the *R* statistical language. In linear modeling, a categorical variable or *factor* can take on two or more values or *levels*. In standard design matrices, one of the values is chosen as a reference value or *base level* and absorbed into the intercept. In standard GLMs, the choice of base level does not influence the values of contrasts (LFCs). This, however, is no longer the case in our approach using ridge-regression-like shrinkage on the coefficients (described below), when factors with more than two levels are present in the design

matrix, because the base level will not undergo shrinkage while the other levels do.

To recover the desirable symmetry between all levels, *DESeq2* uses *expanded design matrices*, which include an indicator variable for *each* level of each factor, in addition to an intercept column (i.e., none of the levels is absorbed into the intercept). While such a design matrix no longer has full rank, a unique solution exists because the zero-centered prior distribution (see below) provides regularization. For dispersion estimation and for estimating the width of the LFC prior, standard design matrices are used.

Contrasts

Contrasts between levels and standard errors of such contrasts can be calculated as they would in the standard design matrix case, i.e., using:

$$\beta_i^c = \tilde{c}^t \tilde{\beta}_i \quad (3)$$

$$\text{SE}(\beta_i^c) = \sqrt{\tilde{c}^t \Sigma_i \tilde{c}}, \quad (4)$$

where $\tilde{c}$ represents a numeric contrast, e.g., 1 and -1 specifying the numerator and denominator of a simple two-level contrast, and $\Sigma_i = \text{Cov}(\tilde{\beta}_i)$, defined below.

Estimation of dispersions

We assume the dispersion parameter α_i follows a log-normal prior distribution that is centered around a trend that depends on the gene's mean normalized read count:

$$\log \alpha_i \sim N(\log \alpha_{\text{tr}}(\bar{\mu}_i), \sigma_d^2). \quad (5)$$

Here, α_{tr} is a function of the gene's mean normalized count,

$$\bar{\mu}_i = \frac{1}{m} \sum_j \frac{K_{ij}}{s_{ij}}.$$

It describes the mean-dependent expectation of the prior. σ_d is the width of the prior, a hyperparameter describing how much the individual genes' true dispersions scatter around the trend. For the trend function, we use the same parametrization as we used for *DEXSeq* [30], namely,

$$\alpha_{\text{tr}}(\bar{\mu}) = \frac{a_1}{\bar{\mu}} + \alpha_0. \quad (6)$$

We get final dispersion estimates from this model in three steps, which implement a computationally fast approximation to a full empirical Bayes treatment. We first use the count data for each gene separately to get preliminary gene-wise dispersion estimates α_i^{gw} by maximum-likelihood estimation. Then, we fit the dispersion trend α_{tr} . Finally, we combine the likelihood with the trended prior to get maximum *a posteriori* (MAP) values as final dispersion estimates. Details for the three steps follow.

Gene-wise dispersion estimates To get a gene-wise dispersion estimate for a gene i , we start by fitting a negative binomial GLM without an LFC prior for the design matrix X to the gene's count data. This GLM uses a rough method-of-moments estimate of dispersion, based on the within-group variances and means. The initial GLM is necessary to obtain an initial set of fitted values, $\hat{\mu}_{ij}^0$. We then maximize the Cox–Reid adjusted likelihood of the dispersion, conditioned on the fitted values $\hat{\mu}_{ij}^0$ from the initial fit, to obtain the gene-wise estimate α_i^{gw} , i.e.,

$$\alpha_i^{\text{gw}} = \arg \max_{\alpha} \ell_{\text{CR}}(\alpha; \bar{\mu}_i^0, \bar{K}_i)$$

with

$$\begin{aligned} \ell_{\text{CR}}(\alpha; \bar{\mu}, \bar{K}) &= \ell(\alpha) - \frac{1}{2} \log(\det(X^t W X)) \\ \ell(\alpha) &= \sum_j \log f_{\text{NB}}(K_j; \mu_j, \alpha), \end{aligned} \quad (7)$$

where $f_{\text{NB}}(k; \mu, \alpha)$ is the probability mass function of the negative binomial distribution with mean μ and dispersion α , and the second term provides the Cox–Reid bias adjustment [47]. This adjustment, first used in the context of dispersion estimation for SAGE data [48] and then for HTS data [3] in *edgeR*, corrects for the negative bias of dispersion estimates from using the MLEs for the fitted values $\hat{\mu}_{ij}^0$ (analogous to Bessel's correction in the usual sample variance formula; for details, see [49], Section 10.6). It is formed from the Fisher information for the fitted values, which is here calculated as $\det(X^t W X)$, where W is the diagonal weight matrix from the standard iteratively reweighted least-squares algorithm. As the GLM's link function is $g(\mu) = \log(\mu)$ and its variance function is $V(\mu; \alpha) = \mu + \alpha\mu^2$, the elements of the diagonal matrix W_i are given by:

$$w_{ji} = \frac{1}{g'(\mu_j)^2 V(\mu_j)} = \frac{1}{1/\mu_j + \alpha}.$$

The optimization in Equation (7) is performed on the scale of $\log \alpha$ using a backtracking line search with accepted proposals that satisfy Armijo conditions [50].

Dispersion trend A parametric curve of the form (6) is fit by regressing the gene-wise dispersion estimates α_i^{gw} onto the means of the normalized counts, $\bar{\mu}_i$. The sampling distribution of the gene-wise dispersion estimate around the true value α_i can be highly skewed, and therefore we do not use ordinary least-squares regression but rather gamma-family GLM regression. Furthermore, dispersion outliers could skew the fit and hence a scheme to exclude such outliers is used.

The hyperparameters α_1 and α_0 of (6) are obtained by iteratively fitting a gamma-family GLM. At each iteration,

genes with a ratio of dispersion to fitted value outside the range $[10^{-4}, 15]$ are left out until the sum of squared LFCs of the new coefficients over the old coefficients is less than 10^{-6} (same approach as in *DEXSeq* [30]).

The parametrization (6) is based on reports by us and others of decreasing dependence of dispersion on the mean in many datasets [3–6,51]. Some caution is warranted to disentangle true underlying dependence from effects of estimation bias that can create a perceived dependence of the dispersion on the mean. Consider a negative binomial distributed random variable with expectation μ and dispersion α . Its variance $v = \mu + \alpha\mu^2$ has two components, $v = v_P + v_D$, the Poisson component $v_P = \mu$ independent of α , and the overdispersion component $v_D = \alpha\mu^2$. When μ is small, $\mu \lesssim 1/\alpha$ (vertical lines in Additional file 1: Figure S1), the Poisson component dominates, in the sense that $v_P/v_D = 1/(\alpha\mu) \gtrsim 1$, and the observed data provide little information on the value of α . Therefore the sampling variance of an estimator for α will be large when $\mu \lesssim 1/\alpha$, which leads to the appearance of bias. For simplicity, we have stated the above argument without regard to the influence of the size factors, s_j , on the value of μ . This is permissible because, by construction, the geometric mean of our size factors is close to 1, and hence, the mean across samples of the unnormalized read counts, $\frac{1}{m} \sum_j K_{ij}$, and the mean of the normalized read counts, $\frac{1}{m} \sum_j K_{ij}/s_j$, will be roughly the same.

This phenomenon may give rise to an apparent dependence of α on μ . It is possible that the shape of the dispersion-mean fit for the Bottomly data (Figure 1A) can be explained in that manner: the asymptotic dispersion is $\alpha_0 \approx 0.01$, and the non-zero slope of the mean-dispersion plot is limited to the range of mean counts up to around 100, the reciprocal of α_0 . However, overestimation of α in that low-count range has little effect on inference, as in that range the variance v is anyway dominated by the α -independent Poisson component v_P . The situation is different for the Pickrell data: here, a dependence of dispersion on mean was observed for counts clearly above the reciprocal of the asymptotic dispersion α_0 (Figure 1B), and hence is not due merely to estimation bias. Simulations (shown in Additional file 1: Figure S25) confirmed that the observed joint distribution of estimated dispersions and means is not compatible with a single, constant dispersion. Therefore, the parametrization (6) is a flexible and mildly conservative modeling choice: it is able to pick up dispersion-mean dependence if it is present, while it can lead to a minor loss of power in the low-count range due to a tendency to overestimate dispersion there.

Dispersion prior As also observed by Wu *et al.* [6], a log-normal prior fits the observed dispersion distribution for typical RNA-seq datasets. We solve the computational

difficulty of working with a non-conjugate prior using the following argument: the logarithmic residuals from the trend fit, $\log \alpha_i^{\text{gw}} - \log \alpha_{\text{tr}}(\bar{\mu}_i)$, arise from two contributions, namely the scatter of the true logarithmic dispersions around the trend, given by the prior with variance σ_d^2 , and the sampling distribution of the logarithm of the dispersion estimator, with variance σ_{de}^2 . The sampling distribution of a dispersion estimator is approximately a scaled χ^2 distribution with $m - p$ degrees of freedom, with m the number of samples and p the number of coefficients. The variance of the logarithm of a χ_f^2 -distributed random variable is given [52] by the trigamma function ψ_1 ,

$$\text{Var} \log X^2 = \psi_1(f/2) \quad \text{for} \quad X^2 \sim \chi_f^2.$$

Therefore, $\sigma_{\text{de}}^2 \approx \psi_1((m - p)/2)$, i.e., the sampling variance of the logarithm of a variance or dispersion estimator is approximately constant across genes and depends only on the degrees of freedom of the model.

Additional file 1: Table S2 compares this approximation for the variance of logarithmic dispersion estimates with the variance of logarithmic Cox–Reid adjusted dispersion estimates for simulated negative binomial data, over a combination of different sample sizes, number of parameters and dispersion values used to create the simulated data. The approximation is close to the sample variance for various typical values of m , p and α .

Therefore, the prior variance σ_d^2 is obtained by subtracting the expected sampling variance from an estimate of the variance of the logarithmic residuals, s_{lr}^2 :

$$\sigma_d^2 = \max \{s_{\text{lr}}^2 - \psi_1((m - p)/2), 0.25\}.$$

The prior variance σ_d^2 is thresholded at a minimal value of 0.25 so that the dispersion estimates are not shrunk entirely to $\alpha_{\text{tr}}(\bar{\mu}_i)$ if the variance of the logarithmic residuals is less than the expected sampling variance.

To avoid inflation of σ_d^2 due to dispersion outliers (i.e., genes not well captured by this prior; see below), we use a robust estimator for the standard deviation s_{lr} of the logarithmic residuals,

$$s_{\text{lr}} = \text{mad}_i (\log \alpha_i^{\text{gw}} - \log \alpha_{\text{tr}}(\bar{\mu}_i)), \quad (8)$$

where mad stands for the median absolute deviation, divided as usual by the scaling factor $\Phi^{-1}(3/4)$.

Three or less residuals degrees of freedom When there are three or less residual degrees of freedom (number of samples minus number of parameters to estimate), the estimation of the prior variance σ_d^2 using the observed variance of logarithmic residuals s_{lr}^2 tends to underestimate σ_d^2 . In this case, we instead estimate the prior variance through simulation. We match the distribution of

logarithmic residuals to a density of simulated logarithmic residuals. These are the logarithm of χ_{m-p}^2 -distributed random variables added to $N(0, \sigma_d^2)$ random variables to account for the spread due to the prior. The simulated distribution is shifted by $-\log(m - p)$ to account for the scaling of the χ^2 distribution. We repeat the simulation over a grid of values for σ_d^2 , and select the value that minimizes the Kullback–Leibler divergence from the observed density of logarithmic residuals to the simulated density.

Final dispersion estimate We form a logarithmic posterior for the dispersion from the Cox–Reid adjusted logarithmic likelihood (7) and the logarithmic prior (5) and use its maximum (i.e., the MAP value) as the final estimate of the dispersion,

$$\alpha_i^{\text{MAP}} = \arg \max_{\alpha} \left(\ell_{\text{CR}} \left(\alpha; \bar{\mu}_i^0, \bar{K}_i \right) + \Lambda_i(\alpha) \right), \quad (9)$$

where

$$\Lambda_i(\alpha) = \frac{-(\log \alpha - \log \alpha_{\text{tr}}(\bar{\mu}_i))^2}{2\sigma_d^2},$$

is, up to an additive constant, the logarithm of the density of prior (5). Again, a backtracking line search is used to perform the optimization.

Dispersion outliers For some genes, the gene-wise estimate α_i^{gw} can be so far above the prior expectation $\alpha_{\text{tr}}(\bar{\mu}_i)$ that it would be unreasonable to assume that the prior is suitable for the gene. If the dispersion estimate for such genes were down-moderated toward the fitted trend, this might lead to false positives. Therefore, we use the heuristic of considering a gene as a dispersion outlier, if the residual from the trend fit is more than two standard deviations of logarithmic residuals, s_{lr} (see Equation (8)), above the fit, i.e., if

$$\log \alpha_i^{\text{gw}} > \log \alpha_{\text{tr}}(\bar{\mu}_i) + 2s_{\text{lr}}.$$

For such genes, the gene-wise estimate α_i^{gw} is not shrunk toward the trended prior mean. Instead of the MAP value α_i^{MAP} , we use the gene-wise estimate α_i^{gw} as a final dispersion value in the subsequent steps. In addition, the iterative fitting procedure for the parametric dispersion trend described above avoids that such dispersion outliers influence the prior mean.

Shrinkage estimation of logarithmic fold changes

To incorporate empirical Bayes shrinkage of LFCs, we postulate a zero-centered normal prior for the coefficients β_{ir} of model (2) that represent LFCs (i.e., typically, all coefficients except for the intercept β_{i0}):

$$\beta_{ir} \sim N(0, \sigma_r^2). \quad (10)$$

As was observed with differential expression analysis using microarrays, genes with low intensity values tend to suffer from a small signal-to-noise ratio. Alternative estimators can be found that are more stable than the standard calculation of fold change as the ratio of average observed values for each condition [53–55]. *DESeq2*'s approach can be seen as an extension of these approaches for stable estimation of gene-expression fold changes to count data.

Empirical prior estimate To obtain values for the empirical prior widths σ_r for the model coefficients, we again approximate a full empirical Bayes approach, as with the estimation of dispersion prior, though here we do not subtract the expected sampling variance from the observed variance of maximum likelihood estimates. The estimate of the LFC prior width is calculated as follows. We use the standard iteratively reweighted least-squares algorithm [12] for each gene's model, Equations (1) and (2), to get MLEs for the coefficients β_{ir}^{MLE} . We then fit, for each column r of the design matrix (except for the intercept), a zero-centered normal distribution to the empirical distribution of MLE fold change estimates β_{ir}^{MLE} .

To make the fit robust against outliers with very high absolute LFC values, we use quantile matching: the width σ_r is chosen such that the $(1-p)$ empirical quantile of the absolute value of the observed LFCs, β_{ir}^{MLE} , matches the $(1-p/2)$ theoretical quantile of the prior, $N(0, \sigma_r^2)$, where p is set by default to 0.05. If we write the theoretical upper quantile of a normal distribution as $Q_N(1-p)$ and the empirical upper quantile of the MLE LFCs as $Q_{|\beta_{ir}|}(1-p)$, then the prior width is calculated as:

$$\sigma_r = \frac{Q_{|\beta_{ir}|}(1-p)}{Q_N(1-p/2)}.$$

To ensure that the prior width σ_r will be independent of the choice of base level, the estimates from the quantile matching procedure are averaged for each factor over all possible contrasts of factor levels. When determining the empirical upper quantile, extreme LFC values ($|\beta_{ir}^{\text{MLE}}| > \log(2) 10$, or 10 on the base 2 scale) are excluded.

Final estimate of logarithmic fold changes The logarithmic posterior for the vector, $\vec{\beta}_i$, of model coefficients β_{ir} for gene i is the sum of the logarithmic likelihood of the GLM (2) and the logarithm of the prior density (10), and its maximum yields the final MAP coefficient estimates:

$$\vec{\beta}_i = \arg \max_{\vec{\beta}} \left(\sum_j \log f_{\text{NB}}(K_{ij}; \mu_j(\vec{\beta}), \alpha_i) + \Lambda(\vec{\beta}) \right),$$

where

$$\mu_j(\vec{\beta}) = s_{ij} e^{\sum_r x_{jr} \beta_r}, \quad \Lambda(\vec{\beta}) = \sum_r \frac{-\beta_r^2}{2\sigma_r^2},$$

and α_i is the final dispersion estimate for gene i , i.e., $\alpha_i = \alpha_i^{\text{MAP}}$, except for dispersion outliers, where $\alpha_i = \alpha_i^{\text{gw}}$.

The term $\Lambda(\vec{\beta})$, i.e., the logarithm of the density of the normal prior (up to an additive constant), can be read as a ridge penalty term, and therefore, we perform the optimization using the *iteratively reweighted ridge regression algorithm* [56], also known as *weighted updates* [57]. Specifically, the updates for a given gene are of the form

$$\vec{\beta} \leftarrow \left(X^t W X + \vec{\lambda} I \right)^{-1} X^t W \vec{z},$$

with $\lambda_r = 1/\sigma_r^2$ and

$$z_j = \log \frac{\mu_j}{s_j} + \frac{K_j - \mu_j}{\mu_j},$$

where the current fitted values $\mu_j = s_{ij} e^{\sum_r x_{jr} \beta_r}$ are computed from the current estimates $\vec{\beta}$ in each iteration.

Fisher information. The effect of the zero-centered normal prior can be understood as shrinking the MAP LFC estimates based on the amount of information the experiment provides for this coefficient, and we briefly elaborate on this here. Specifically, for a given gene i , the shrinkage for an LFC β_{ir} depends on the *observed Fisher information*, given by

$$\mathcal{I}_m(\hat{\beta}_{ir}) = - \left[\frac{\partial^2}{\partial \beta_{ir}^2} \ell(\vec{\beta}_i; \vec{K}_i, \alpha_i) \right]_{\beta_{ir} = \hat{\beta}_{ir}},$$

where $\ell(\vec{\beta}_i; \vec{K}_i, \alpha_i)$ is the logarithm of the likelihood, and partial derivatives are taken with respect to LFC β_{ir} . For a negative binomial GLM, the observed Fisher information, or peakedness of the logarithm of the profile likelihood, is influenced by a number of factors including the degrees of freedom, the estimated mean counts μ_{ij} , and the gene's dispersion estimate α_i . The prior influences the MAP estimate when the density of the likelihood and the prior are multiplied to calculate the posterior. Genes with low estimated mean values μ_{ij} or high dispersion estimates α_i have flatter profile likelihoods, as do datasets with few residual degrees of freedom, and therefore in these cases the zero-centered prior pulls the MAP estimate from a high-uncertainty MLE closer toward zero.

Wald test

The Wald test compares the beta estimate β_{ir} divided by its estimated standard error $\text{SE}(\beta_{ir})$ to a standard normal distribution. The estimated standard errors are the square root of the diagonal elements of the estimated covariance matrix, Σ_i , for the coefficients, i.e., $\text{SE}(\beta_{ir}) = \sqrt{\Sigma_{i,rr}}$. Contrasts of coefficients are tested similarly by forming a Wald statistics using (3) and (4). We use the following

formula for the coefficient covariance matrix for a GLM with normal prior on coefficients [56,58]:

$$\Sigma_i = \text{Cov}(\vec{\beta}_i) = (X^t W X + \vec{\lambda} I)^{-1} (X^t W X) (X^t W X + \vec{\lambda} I)^{-1}.$$

The tail integrals of the standard normal distribution are multiplied by 2 to achieve a two-tailed test. The Wald test P values from the subset of genes that pass the independent filtering step are adjusted for multiple testing using the procedure of Benjamini and Hochberg [21].

Independent filtering

Independent filtering does not compromise type-I error control as long as the distribution of the test statistic is marginally independent of the filter statistic *under the null hypothesis* [22], and we argue in the following that this is the case in our application. The filter statistic in *DESeq2* is the mean of normalized counts for a gene, while the test statistic is p , the P value from the Wald test. We first consider the case where the size factors are equal and where the gene-wise dispersion estimates are used for each gene, i.e. without dispersion shrinkage. The distribution family for the negative binomial is parameterized by $\theta = (\mu, \alpha)$. Aside from discreteness of p due to low counts, for a given μ , the distribution of p is Uniform(0, 1) under the null hypothesis, so p is an ancillary statistic. The sample mean of counts for gene i , $\bar{K}_i$, is boundedly complete sufficient for μ . Then from Basu's theorem, $\bar{K}_i$ and p are independent.

While for very low counts, one can observe discreteness and non-uniformity of p under the null hypothesis, *DESeq2* does not use the distribution of p in its estimation procedure – for example, *DESeq2* does not estimate the proportion of null genes using the distribution of p – so this kind of dependence of p on μ does not lead to increased type-I error.

If the size factors are not equal across samples, but not correlated with condition, conditioning on the mean of *normalized* counts should also provide uniformly distributed p as with conditioning on the mean of counts, $\bar{K}_i$. We may consider a pathological case where the size factors are perfectly confounded with condition, in which case, even under the null hypothesis, genes with low mean count would have non-uniform distribution of p , as one condition could have positive counts and the other condition often zero counts. This could lead to non-uniformity of p under the null hypothesis; however, such a pathological case would pose problems for many statistical tests of differences in mean.

We used simulation to demonstrate that the independence of the null distribution of the test statistic from the filter statistic still holds for dispersion shrinkage. Additional file 1: Figure S26 displays marginal null distributions of p across the range of mean normalized counts. Despite spikes in the distribution for the genes with the

lowest mean counts due to discreteness of the data, these densities were nearly uniform across the range of average expression strength.

Composite null hypotheses

DESeq2 offers tests for composite null hypotheses of the form $\mathcal{H}_0 : |\beta_{ir}| \leq \theta$ to find genes whose LFC significantly exceeds a threshold $\theta > 0$. The composite null hypothesis is replaced by two simple null hypotheses: $\mathcal{H}_{0a} : \beta_{ir} = \theta$ and $\mathcal{H}_{0b} : \beta_{ir} = -\theta$. Two-tailed P values are generated by integrating a normal distribution centered on θ with standard deviation $\text{SE}(\beta_{ir})$ from $|\beta_{ir}|$ toward ∞ . The value of the integral is then multiplied by 2 and thresholded at 1. This procedure controls type-I error even when $\beta_{ir} = \pm\theta$, and is equivalent to the standard *DESeq2* P value when $\theta = 0$.

Conversely, when searching for genes whose absolute LFC is significantly below a threshold, i.e., when testing the null hypothesis $\mathcal{H}_0 : |\beta_{ir}| \geq \theta$, the P value is constructed as the maximum of two one-sided tests of the simple null hypotheses: $\mathcal{H}_{0a} : \beta_{ir} = \theta$ and $\mathcal{H}_{0b} : \beta_{ir} = -\theta$. The one-sided P values are generated by integrating a normal distribution centered on θ with standard deviation $\text{SE}(\beta_{ir})$ from β_{ir} toward $-\infty$, and integrating a normal distribution centered on $-\theta$ with standard deviation $\text{SE}(\beta_{ir})$ from β_{ir} toward ∞ .

Note that while a zero-centered prior on LFCs is consistent with testing the null hypothesis of small LFCs, it should not be used when testing the null hypothesis of large LFCs, because the prior would then favor the alternative hypothesis. *DESeq2* requires that no prior has been used when testing the null hypothesis of large LFCs, so that the data alone must provide evidence against the null hypothesis.

Interactions

Two exceptions to the default *DESeq2* LFC estimation steps are used for experimental designs with interaction terms. First, when any interaction terms are included in the design, the LFC prior width for main effect terms is not estimated from the data, but set to a wide value ($\sigma_r^2 = (\log(2))^2 1000$, or 1000 on the base 2 scale). This ensures that shrinkage of main effect terms will not result in false positive calls of significance for interactions. Second, when interaction terms are included and all factors have two levels, then standard design matrices are used rather than expanded model matrices, such that only a single term is used to test the null hypothesis that a combination of two effects is merely additive in the logarithmic scale.

Regularized logarithm

The rlog transformation is calculated as follows. The experimental design matrix X is substituted with a design

matrix with an indicator variable for every sample in addition to an intercept column. A model as described in Equations (1) and (2) is fit with a zero-centered normal prior on the non-intercept terms and using the fitted dispersion values $\alpha_{tr}(\bar{\mu})$, which capture the overall variance-mean dependence of the dataset. The true experimental design matrix X is then only used in estimating the variance-mean trend over all genes. For unsupervised analyses, for instance sample quality assessment, it is desirable that the experimental design has no influence on the transformation, and hence *DESeq2* by default ignores the design matrix and re-estimates the dispersions treating all samples as replicates, i.e., it uses *blind* dispersion estimation. The rlog-transformed values are the fitted values,

$$\text{rlog}(K_{ij}) \equiv \log_2 q_{ij} = \beta_{i0} + \beta_{ij},$$

where β_{ij} is the shrunken LFC on the base 2 scale for the j th sample. The variance of the prior is set using a similar approach as taken with differential expression, by matching a zero-centered normal distribution to observed LFCs. First a matrix of LFCs is calculated by taking the logarithm (base 2) of the normalized counts plus a pseudocount of $\frac{1}{2}$ for each sample divided by the mean of normalized counts plus a pseudocount of $\frac{1}{2}$. The pseudocount of $\frac{1}{2}$ allows for calculation of the logarithmic ratio for all genes, and has little effect on the estimate of the variance of the prior or the final rlog transformation. This matrix of LFCs then represents the common-scale logarithmic ratio of each sample to the fitted value using only an intercept. The prior variance is found by matching the 97.5% quantile of a zero-centered normal distribution to the 95% quantile of the absolute values in the LFC matrix.

Cook's distance for outlier detection

The MLE of $\bar{\beta}_i$ is used for calculating Cook's distance. Considering a gene i and sample j , Cook's distance for GLMs is given by [59]:

$$D_{ij} = \frac{R_{ij}^2}{\tau p} \frac{h_{jj}}{(1 - h_{jj})^2},$$

where R_{ij} is the Pearson residual of sample j , τ is an overdispersion parameter (in the negative binomial GLM, τ is set to 1), p is the number of parameters including the intercept, and h_{jj} is the j th diagonal element of the hat matrix H :

$$H = W^{1/2} X (X^t W X)^{-1} X^t W^{1/2}.$$

Pearson residuals R_{ij} are calculated as

$$R_{ij} = \frac{(K_{ij} - \mu_{ij})}{\sqrt{V(\mu_{ij})}},$$

where μ_{ij} is estimated by the negative binomial GLM without the LFC prior, and using the variance function $V(\mu) = \mu + \alpha\mu^2$. A method-of-moments estimate α_i^{rob} , using a robust estimator of variance $s_{i,\text{rob}}^2$ to provide robustness against outliers, is used here:

$$\alpha_i^{\text{rob}} = \max \left(\frac{s_{i,\text{rob}}^2 - \bar{\mu}_i}{\bar{\mu}_i^2}, 0 \right).$$

R/Bioconductor package

DESeq2 is implemented as a package for the R statistical environment and is available [10] as part of the Bioconductor project [11]. The count matrix and metadata, including the gene model and sample information, are stored in an S4 class derived from the *SummarizedExperiment* class of the *GenomicRanges* package [60]. *SummarizedExperiment* objects containing count matrices can be easily generated using the *summarizeOverlaps* function of the *GenomicAlignments* package [61]. This workflow automatically stores the gene model as metadata and additionally other information such as the genome and gene annotation versions. Other methods to obtain count matrices include the *htseq-count* script [62] and the Bioconductor packages *easyRNASeq* [63] and *feature-Count* [64].

The *DESeq2* package comes with a detailed vignette, which works through a number of example differential expression analyses on real datasets, and the use of the rlog transformation for quality assessment and visualization. A single function, called *DESeq*, is used to run the default analysis, while lower-level functions are also available for advanced users.

Read alignment for the Bottomly et al. and Pickrell et al. datasets

Reads were aligned using the TopHat2 aligner [65], and assigned to genes using the *summarizeOverlaps* function of the *GenomicRanges* package [60]. The sequence read archive fastq files of the Pickrell et al. [17] dataset (accession number [SRA:SRP001540]) were aligned to the *Homo sapiens* reference sequence GRCh37 downloaded in March 2013 from Illumina iGenomes. Reads were counted in the genes defined by the Ensembl GTF file, release 70, contained in the Illumina iGenome. The sequence read archive fastq files of the Bottomly et al. [16] dataset (accession number [SRA:SRP004777]) were aligned to the *Mus musculus* reference sequence NCBI37 downloaded in March 2013 from Illumina iGenomes. Reads were counted in the genes defined by

the Ensembl GTF file, release 66, contained in the Illumina iGenome.

Reproducible code

Sweave vignettes for reproducing all figures and tables in this paper, including data objects for the experiments mentioned, and code for aligning reads and for benchmarking, can be found in a package *DESeq2paper* [66].

Additional file

Additional file 1: Supplementary methods, tables and figures.

Abbreviations

FDR: False discovery rate; GLM: Generalized linear model; HTS: High-throughput sequencing; LFC: Logarithmic fold change; MAP: Maximum a posteriori; MLE: Maximum-likelihood estimate; RNA-seq: RNA sequencing; VST: Variance-stabilizing transformation.

Competing interests

The authors declare that they have no competing interests.

Authors' contributions

All authors developed the method and wrote the manuscript. MIL implemented the method and performed the analyses. All authors read and approved the final manuscript.

Acknowledgements

The authors thank all users of DESeq and DESeq2 who provided valuable feedback. We thank Judith Zaugg for helpful comments on the manuscript. MIL acknowledges funding via a stipend from the International Max Planck Research School for Computational Biology and Scientific Computing and a grant from the National Institutes of Health (5T32CA009337-33). WH and SA acknowledge funding from the European Union's 7th Framework Programme (Health) via Project *Radiant*. We thank an anonymous reviewer for raising the question of estimation biases in the dispersion-mean trend fitting.

Author details

¹Department of Biostatistics and Computational Biology, Dana Farber Cancer Institute and Department of Biostatistics, Harvard School of Public Health, 450 Brookline Avenue, Boston, MA 02215, USA. ²Genome Biology Unit, European Molecular Biology Laboratory, Meyerhofstrasse 1, 69117 Heidelberg, Germany.

³Department of Computational Molecular Biology, Max Planck Institute for Molecular Genetics, Ihnestrasse 63-73, 14195 Berlin, Germany.

Received: 27 May 2014 Accepted: 19 November 2014

Published online: 05 December 2014

References

- Lönnstedt I, Speed T: **Replicated microarray data.** *Stat Sinica* 2002, **12**:31–46.
- Robinson MD, Smyth GK: **Moderated statistical tests for assessing differences in tag abundance.** *Bioinformatics* 2007, **23**:2881–2887.
- McCarthy DJ, Chen Y, Smyth GK: **Differential expression analysis of multifactor RNA-seq experiments with respect to biological variation.** *Nucleic Acids Res* 2012, **40**:4288–4297.
- Anders S, Huber W: **Differential expression analysis for sequence count data.** *Genome Biol* 2010, **11**:106.
- Zhou Y-H, Xia K, Wright FA: **A powerful and flexible approach to the analysis of RNA sequence count data.** *Bioinformatics* 2011, **27**:2672–2678.
- Wu H, Wang C, Wu Z: **A new shrinkage estimator for dispersion improves differential expression detection in RNA-seq data.** *Biostatistics* 2013, **14**:232–243.
- Hardcastle T, Kelly K: **baySeq: empirical Bayesian methods for identifying differential expression in sequence count data.** *BMC Bioinformatics* 2010, **11**:422.
- Van De Wiel MA, Leday GGR, Pardo L, Rue H, Van Der Vaart AW, Van Wieringen WN: **Bayesian analysis of RNA sequencing data by estimating multiple shrinkage priors.** *Biostatistics* 2013, **14**:113–128.
- Boer JM, Huber WK, Sülthmann H, Wilmer F, von Heydebreck A, Haas S, Korn B, Gunawan B, Vente A, Füzesi L, Vingron M, Poustka A: **Identification and classification of differentially expressed genes in renal cell carcinoma by expression profiling on a global human 31,500-element cDNA array.** *Genome Res* 2001, **11**:1861–1870.
- DESeq2.** [http://www.bioconductor.org/packages/release/bioc/html/DESeq2.html]
- Gentleman RC, Carey VJ, Bates DM, Bolstad B, Dettling M, Dudoit S, Ellis B, Gautier L, Ge Y, Gentry J, Hornik K, Hothorn T, Huber W, Iacus S, Irizarry R, Leisch F, Li C, Maechler M, Rossini AJ, Sawitzki G, Smith C, Smyth G, Tierney L, Yang JY, Zhang J: **Bioconductor: open software development for computational biology and bioinformatics.** *Genome Biol* 2004, **5**:R80.
- McCullagh P, Nelder JA: **Generalized linear models.** In *Monographs on Statistics & Applied Probability*. 2nd edition. London, UK: Chapman & Hall/CRC; 1989.
- Hansen KD, Irizarry RA, Wu Z: **Removing technical variability in RNA-seq data using conditional quantile normalization.** *Biostatistics* 2012, **13**:204–216.
- Risso D, Schwartz K, Sherlock G, Dudoit S: **GC-content normalization for RNA-seq data.** *BMC Bioinformatics* 2011, **12**:480.
- Smyth GK: **Linear models and empirical Bayes methods for assessing differential expression in microarray experiments.** *Stat Appl Genet Mol Biol* 2004, **3**:1–25.
- Bottomly D, Walter NAR, Hunter JE, Darakjian P, Kawane S, Buck KJ, Searles RP, Mooney M, McWeeney SK, Hitzemann R: **Evaluating gene expression in C57BL/6J and DBA/2J mouse striatum using RNA-seq and microarrays.** *PLoS ONE* 2011, **6**:17820.
- Pickrell JK, Marioni JC, Pai AA, Degner JF, Engelhardt BE, Nkadori E, Veyrieras J-B, Stephens M, Gilad Y, Pritchard JK: **Understanding mechanisms underlying human gene expression variation with RNA sequencing.** *Nature* 2010, **464**:768–772.
- Hastie T, Tibshirani R, Friedman J: *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*. 2nd edition. New York City, USA: Springer; 2009.
- Bi Y, Davuluri R: **NPEBseq: nonparametric empirical Bayesian-based procedure for differential expression analysis of RNA-seq data.** *BMC Bioinformatics* 2013, **14**:262.
- Feng J, Meyer CA, Wang Q, Liu XS, Zhang Y: **GFOLD: a generalized fold change for ranking differentially expressed genes from RNA-seq data.** *Bioinformatics* 2012, **28**:2782–2788.
- Benjamini Y, Hochberg Y: **Controlling the false discovery rate: a practical and powerful approach to multiple testing.** *J R Stat Soc Ser B Methodol* 1995, **57**:289–300.
- Bourgon R, Gentleman R, Huber W: **Independent filtering increases detection power for high-throughput experiments.** *Proc Natl Acad Sci USA* 2010, **107**:9546–9551.
- McCarthy DJ, Smyth GK: **Testing significance relative to a fold-change threshold is a TREAT.** *Bioinformatics* 2009, **25**:765–771.
- Li J, Tibshirani R: **Finding consistent patterns: a nonparametric approach for identifying differential expression in RNA-seq data.** *Stat Methods Med Res* 2013, **22**:519–536.
- Cook RD: **Detection of influential observation in linear regression.** *Technometrics* 1977, **19**:15–18.
- Hammer P, Banck MS, Amberg R, Wang C, Petznick G, Luo S, Khrebukova I, Schroth GP, Beyerlein P, Beutler AS: **mRNA-seq with agnostic splice site discovery for nervous system transcriptomics tested in chronic pain.** *Genome Res* 2010, **20**:847–860.
- Frazee A, Langmead B, Leek J: **ReCount: a multi-experiment resource of analysis-ready RNA-seq gene count datasets.** *BMC Bioinformatics* 2011, **12**:449.
- Trapnell C, Hendrickson DG, Sauvageau M, Goff L, Rinn JL, Pachter L: **Differential analysis of gene regulation at transcript resolution with RNA-seq.** *Nat Biotechnol* 2012, **31**:46–53.
- Glaus P, Honkela A, Rattray M: **Identifying differentially expressed transcripts from RNA-seq data with biological variation.** *Bioinformatics* 2012, **28**:1721–1728.

30. Anders S, Reyes A, Huber W: **Detecting differential usage of exons from RNA-seq data.** *Genome Res* 2012, **22**:2008–2017.
31. Sammeth M: **Complete alternative splicing events are bubbles in splicing graphs.** *J Comput Biol* 2009, **16**:1117–1140.
32. Pagès H, Bindreither D, Carlson M, Morgan M: **SplicingGraphs: create, manipulate, visualize splicing graphs, and assign RNA-seq reads to them.** 2013. Bioconductor package [http://www.bioconductor.org]
33. Robinson MD, McCarthy DJ, Smyth GK: **edgeR: a Bioconductor package for differential expression analysis of digital gene expression data.** *Bioinformatics* 2009, **26**:139–140.
34. Zhou X, Lindsay H, Robinson MD: **Robustly detecting differential expression in RNA sequencing data using observation weights.** *Nucleic Acids Res* 2014, **42**:e91.
35. Leng N, Dawson JA, Thomson JA, Ruotti V, Rissman AI, Smits BMG, Haag JD, Gould MN, Stewart RM, Kendziorski C: **EBSeq: an empirical Bayes hierarchical model for inference in RNA-seq experiments.** *Bioinformatics* 2013, **29**:1035–1043.
36. Law CW, Chen Y, Shi W, Smyth GK: **Voom: precision weights unlock linear model analysis tools for RNA-seq read counts.** *Genome Biol* 2014, **15**:29.
37. Hubert L, Arabie P: **Comparing partitions.** *J Classif* 1985, **2**:193–218.
38. Witten DM: **Classification and clustering of sequencing data using a Poisson model.** *Ann Appl Stat* 2011, **5**:2493–2518.
39. Irizarry RA, Wu Z, Jaffee HA: **Comparison of affymetrix GeneChip expression measures.** *Bioinformatics* 2006, **22**:789–794.
40. Asangani IA, Dommeti VL, Wang X, Malik R, Cieslik M, Yang R, Escara-Wilke J, Wilder-Romans K, Dhanireddy S, Engelke C, Iyer MK, Jing X, Wu Y-M, Cao X, Qin ZS, Wang S, Feng FY, Chinnaiyan AM: **Therapeutic targeting of BET bromodomain proteins in castration-resistant prostate cancer.** *Nature* 2014, **510**:278–282.
41. Stark R, Brown G: **DiffBind: differential binding analysis of ChIP-seq peak data.** 2013. Bioconductor package [http://www.bioconductor.org]
42. Ross-Innes CS, Stark R, Teschendorff AE, Holmes KA, Ali HR, Dunning MJ, Brown GD, Gojis O, Ellis IO, Green AR, Ali S, Chin S-F, Palmieri C, Caldas C, Carroll JS: **Differential oestrogen receptor binding is associated with clinical outcome in breast cancer.** *Nature* 2012, **481**:389–393.
43. Robinson DG, Chen W, Storey JD, Gresham D: **Design and analysis of bar-seq experiments.** *G3 (Bethesda)* 2013, **4**:11–18.
44. McMurdie PJ, Holmes S: **Waste not, want not: why rarefying microbiome data is inadmissible.** *PLoS Comput Biol* 2014, **10**:1003531.
45. Vasquez J, Hon C, Vanselow JT, Schlosser A, Siegel TN: **Comparative ribosome profiling reveals extensive translational complexity in different *Trypanosoma brucei* life cycle stages.** *Nucleic Acids Res* 2014, **42**:3623–3637.
46. Zhou Y, Zhu S, Cai C, Yuan P, Li C, Huang Y, Wei W: **High-throughput screening of a CRISPR/Cas9 library for functional genomics in human cells.** *Nature* 2014, **509**:487–491.
47. Cox DR, Reid N: **Parameter orthogonality and approximate conditional inference.** *J R Stat Soc Ser B Methodol* 1987, **49**:1–39.
48. Robinson MD, Smyth GK: **Small-sample estimation of negative binomial dispersion, with applications to SAGE data.** *Biostatistics* 2007, **9**:321–332.
49. Pawitan Y: *In All Likelihood: Statistical Modelling and Inference Using Likelihood.* 1st edition. New York City, USA: Oxford University Press; 2001.
50. Armijo L: **Minimization of functions having Lipschitz continuous first partial derivatives.** *Pac J Math* 1966, **16**:1–3.
51. Di Y, Schafer DW, Cumbie JS, Chang JH: **The NBP negative binomial model for assessing differential gene expression from RNA-seq.** *Stat Appl Genet Mol Biol* 2011, **10**:1–28.
52. Abramowitz M, Stegun I: *Handbook of Mathematical Functions.* 1st edition. New York, USA: Dover Publications; 1965. [Dover Books on mathematics]
53. Newton M, Kendziorski C, Richmond C, Blattner F, Tsui K: **On differential variability of expression ratios: improving statistical inference about gene expression changes from microarray data.** *J Comput Biol* 2001, **8**:37–52.
54. Huber W, von Heydebreck A, Sultmann H, Poustka A, Vingron M: **Variance stabilization applied to microarray data calibration and to the quantification of differential expression.** *Bioinformatics* 2002, **18**:96–104.
55. Durbin BP, Hardin JS, Hawkins DM, Rocke DM: **A variance-stabilizing transformation for gene-expression microarray data.** *Bioinformatics* 2002, **18**:105–110.
56. Park MY: **Generalized linear models with regularization.** *PhD thesis.* Stanford University, Department of Statistics; 2006.
57. Friedman J, Hastie T, Tibshirani R: **Regularization paths for generalized linear models via coordinate descent.** *J Stat Softw* 2010, **33**:1–22.
58. Cule E, Vineis P, De Iorio M: **Significance testing in ridge regression for genetic data.** *BMC Bioinformatics* 2011, **12**:372.
59. Cook RD, Weisberg S: *Residuals and Influence in Regression.* 1st edition. New York, USA: Chapman and Hall/CRC; 1982.
60. Lawrence M, Huber W, Pagès H, Aboyoun P, Carlson M, Gentleman R, Morgan MT, Carey VJ: **Software for computing and annotating genomic ranges.** *PLoS Comput Biol* 2013, **9**:1003118.
61. Pagès H, Obenchain V, Morgan M: **GenomicAlignments: Representation and manipulation of short genomic alignments.** 2013. Bioconductor package [http://www.bioconductor.org]
62. Anders S, Pyl PT, Huber W: **HTSeq - A Python framework to work with high-throughput sequencing data.** *Bioinformatics* 2015, **31**:166.
63. Delhomme N, Padioulet I, Furlong EE, Steinmetz LM: **easyRNASeq: a Bioconductor package for processing RNA-seq data.** *Bioinformatics* 2012, **28**:2532–2533.
64. Liao Y, Smyth GK, Shi W: **featureCounts: an efficient general purpose program for assigning sequence reads to genomic features.** *Bioinformatics* 2014, **30**:923–930.
65. Kim D, Pertea G, Trapnell C, Pimentel H, Kelley R, Salzberg S: **TopHat2: accurate alignment of transcriptomes in the presence of insertions, deletions and gene fusions.** *Genome Biol* 2013, **14**:36.
66. **DESeq2paper.** [http://www-huber.embl.de/DESeq2paper]

Submit your next manuscript to BioMed Central and take full advantage of:

- **Convenient online submission**
- **Thorough peer review**
- **No space constraints or color figure charges**
- **Immediate publication on acceptance**
- **Inclusion in PubMed, CAS, Scopus and Google Scholar**
- **Research which is freely available for redistribution**

Submit your manuscript at
www.biomedcentral.com/submit

