

REVIEW

Open Access



From bulk RNA sequencing to spatial transcriptomics: a comparative review of differential gene expression analysis methods

Maryam Kalantari-Dehaghi¹ , Neda Ghohabi-Esfahani^{1,2} and Modjtaba Emadi-Baygi^{1*}

Abstract

Background Transcriptome analysis is essential to dissect the molecular mechanisms underlying phenotypic differences and disease mechanisms. Traditionally, microarrays were used, but RNA sequencing (RNA-seq) has become a more powerful and reproducible alternative. RNA-seq is widely applied for differential expression analysis (DEA) analysis, making it possible to characterize gene expression heterogeneity across diverse biological conditions.

Main body Several approaches are used in RNA-seq, including bulk RNA-seq, single-cell RNA sequencing (scRNA-seq), direct RNA sequencing (DRS), and spatial transcriptomics (ST), every distinct offering unique insight. Bulk RNA-seq offers a global perspective on transcriptional activity, while scRNA-seq dissects the intrinsic cellular diversity and developmental trajectories, aiding for delineating molecular profiles of low-abundance cell phenotypes and complex biology. However, both methods face challenges such as biases from short-read sequencing, limited isoform resolution, and the loss of spatial information. Advanced methods such as DRS and ST have emerged to address these limitations. RNA-seq data analysis comprises a multi-phase workflow, including read trimming, alignment, quantification, normalization, and differential expression testing. As sequencing technologies evolve, numerous computational tools and statistical methods have been developed to support DEA across bulk, scRNA-seq, DRS, and ST datasets, enhancing our ability to interpret transcriptomic changes in health and disease. Applications of RNA-seq include interrogation of multifactorial disease mechanisms, biomarker discovery, and the investigation of gene regulatory networks.

Conclusion Examining gene expression through these diverse RNA-seq modalities deepens our understanding of cellular processes, transcript modifications, tissue spatial organization, and gene expression variability. This review highlights current DEA methods in R and Python for bulk RNA-seq, scRNA-seq, DRS, and ST, showcasing the rapidly evolving landscape of transcriptome research.

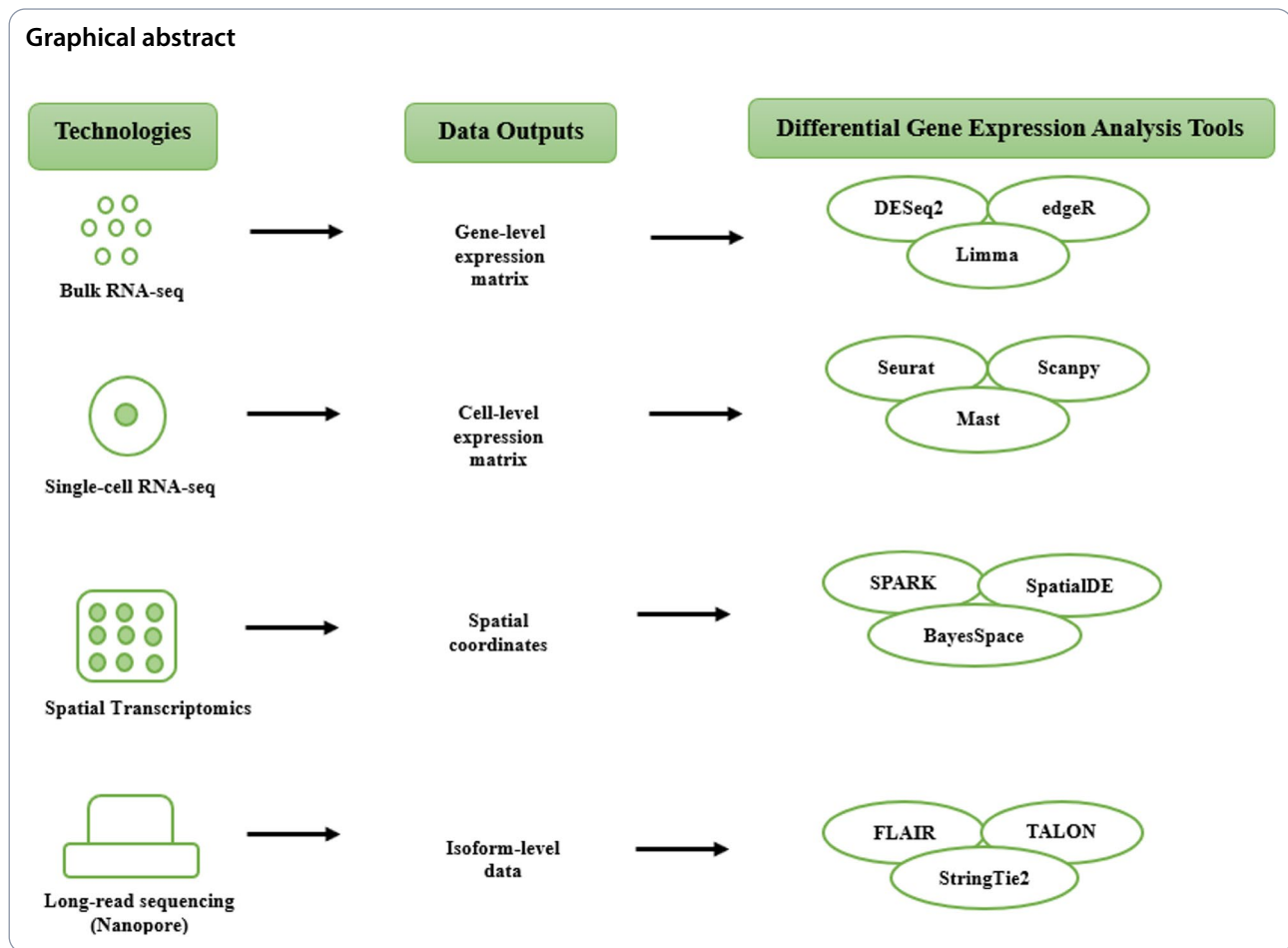
Keywords Bulk RNA-seq, scRNA-seq, Direct RNA sequencing, Spatial transcriptomics, DEA, R, Python

*Correspondence:
Modjtaba Emadi-Baygi
emadi-m@sku.ac.ir; emadibaygi@gmail.com

Full list of author information is available at the end of the article



© The Author(s) 2025. **Open Access** This article is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License, which permits any non-commercial use, sharing, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if you modified the licensed material. You do not have permission under this licence to share adapted material derived from this article or parts of it. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by-nc-nd/4.0/>.

Graphical abstract**Background**

Understanding how genes drive traits and disease relies heavily on studying the transcriptome, the full set of RNA molecules in a cell or tissue. For many years, microarrays were the main tool for this kind of analysis. However, with the rise of high-throughput sequencing, RNA sequencing (RNA-seq) has become the standard approach [1, 2]. RNA-seq allows researchers to measure how much each gene or transcript is expressed, enabling Differential Expression Analysis (DEA) to identify which genes change significantly between different conditions. Compared to microarrays, RNA-seq is more sensitive, provides higher resolution, and produces more reproducible results [3–6]. A typical RNA-seq experiment involves several key steps: preprocessing the sequencing reads, aligning them to a reference genome, quantifying expression levels, normalizing the data, and performing DEA. Beyond identifying differentially expressed genes (DEGs), RNA-seq can be used for many other purposes, such as discovering biomarkers, mapping gene networks, and analyzing biological pathways [7]. Today, RNA-seq comes in several flavors, including bulk RNA sequencing, single-cell RNA sequencing (scRNA-seq), direct RNA

sequencing (DRS), and spatial transcriptomics (ST) (Fig. 1).

Bulk RNA-seq measures the average gene expression across a large population of cells, while single-cell RNA sequencing (scRNA-seq) looks at each cell individually. This single-cell approach allows researchers to uncover cellular heterogeneity, identify rare cell types, and track developmental trajectories or “pseudotime” dynamics [8–13]. While DEA is still an important part of scRNA-seq studies, the main power of this technology lies in exploring the diversity of cell types and dynamic cellular states. It’s worth noting that applying traditional bulk RNA-seq methods directly to single-cell data can lead to problems, due to issues like sparse data, dropout events, and overdispersion [14–16]. To address these challenges, specialized statistical models have been developed to handle the unique features of single-cell datasets [17, 18]. A variety of computational tools are now available for DEA and related analyses. Many are implemented in R, such as DESeq2, edgeR, and limma, which remain widely used. At the same time, Python-based frameworks like Scanpy, PyDESeq2, and Scedar are gaining popularity,

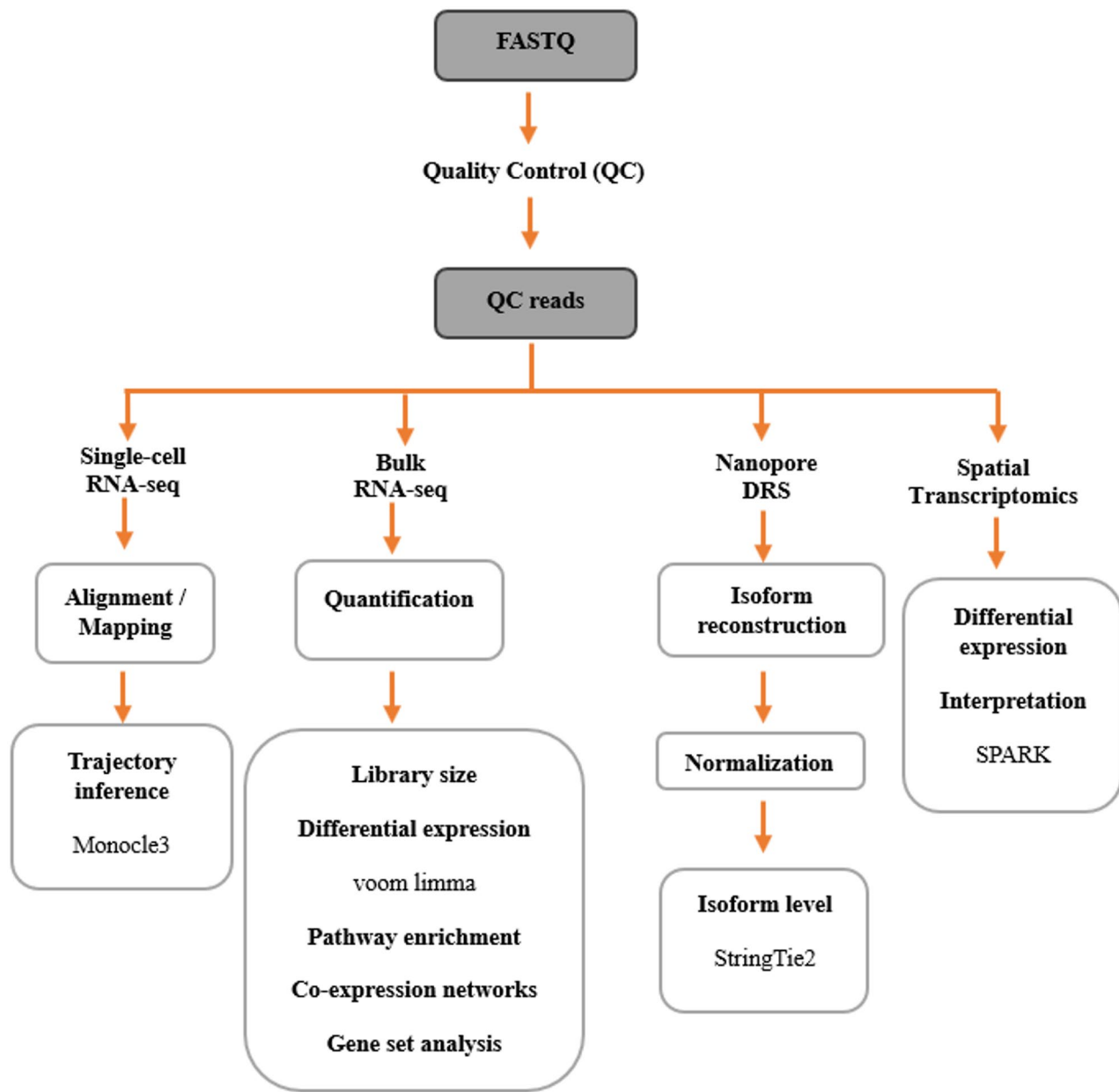


Fig. 1 Schematic of the RNA-seq workflow, including read preprocessing, alignment, quantification, and downstream analyses such as normalization and differential expression. Bold labels indicate pipeline stages; tool names are shown in regular font

reflecting a broader trend toward flexible, interoperable bioinformatics pipelines [19, 20].

Advances in sequencing technology have expanded what RNA sequencing can do. With new long-read platforms like Oxford Nanopore Technologies (ONT), scientists can now sequence entire RNA molecules directly, without first converting them to DNA. This means they can see the original RNA modifications and capture all transcript variants more accurately. This approach is especially helpful for measuring different transcript isoforms and detecting RNA modifications, rather than just comparing gene expression levels like traditional methods [21–25]. At the same time, spatial transcriptomics allows researchers to see where genes are active within intact tissue sections. Unlike bulk or single-cell RNA-seq, which lose the spatial location of cells, this method maps gene expression while preserving tissue structure. This lets scientists’ study how cells interact, their arrangement, and how their environment affects gene activity. Analyses here look for genes that vary by location instead of just overall expression differences, offering new insights into tissue complexity in areas like cancer, brain research, and development [26–29].

Across all types of RNA sequencing, it remains crucial to carefully normalize the data, apply appropriate statistical models, and use visualization techniques to ensure that DEA results are accurate and can be replicated. Recently, integrative frameworks that combine bulk RNA-seq, single-cell RNA-seq, and spatial transcriptomics data have started to provide a broader, multi-level understanding of how tissues are organized and how gene regulation operates. This marks an exciting direction for future transcriptomic research. Collectively, these various technologies, bulk, single-cell, direct RNA, and spatial transcriptomics, offer complementary perspectives on gene expression across different levels of biological complexity. In this review, computational methods for DEA and related studies across these RNA-seq types are summarized, focusing especially on tools available in R and Python (see Fig. 2; Table 1). It's also important to note that RNA-seq methods for DEA have some limitations, which are outlined in Table 2.

Overview of differential expression analysis methods

DEA plays a central role in making sense of transcriptomic data across a range of sequencing technologies, including bulk RNA-seq, scRNA-seq, DRS, and ST. Each approach comes with its own strengths and analytical

challenges, meaning that specialized computational methods are needed to accurately measure and interpret gene expression. Bulk RNA-seq remains the standard method for identifying DEGs across different biological conditions. It provides an averaged view of transcription across large populations of cells, making it useful for broad comparisons. In contrast, scRNA-seq offers a much more detailed picture by profiling gene expression at the single-cell level. This approach can uncover cellular heterogeneity, lineage relationships, and dynamic transitions. DEA in scRNA-seq is usually applied to compare clusters, cell types, or conditions, while accounting for the inherent sparsity and variability of single-cell data. DRS takes transcriptomics a step further by sequencing native RNA molecules directly, without the need for reverse transcription or amplification. This allows researchers to detect RNA modifications and full-length isoforms, providing a deeper understanding of transcript diversity and post-transcriptional regulation. While DRS data can be used to quantify gene expression, its primary strength lies in isoform analysis and mapping RNA modifications, rather than classic DEG identification. ST adds yet another layer by combining gene expression profiling with spatial information within tissue sections. This technology allows scientists to explore how gene activity is influenced by cell–cell interactions and tissue micro-environments. Analytical approaches in ST often focus

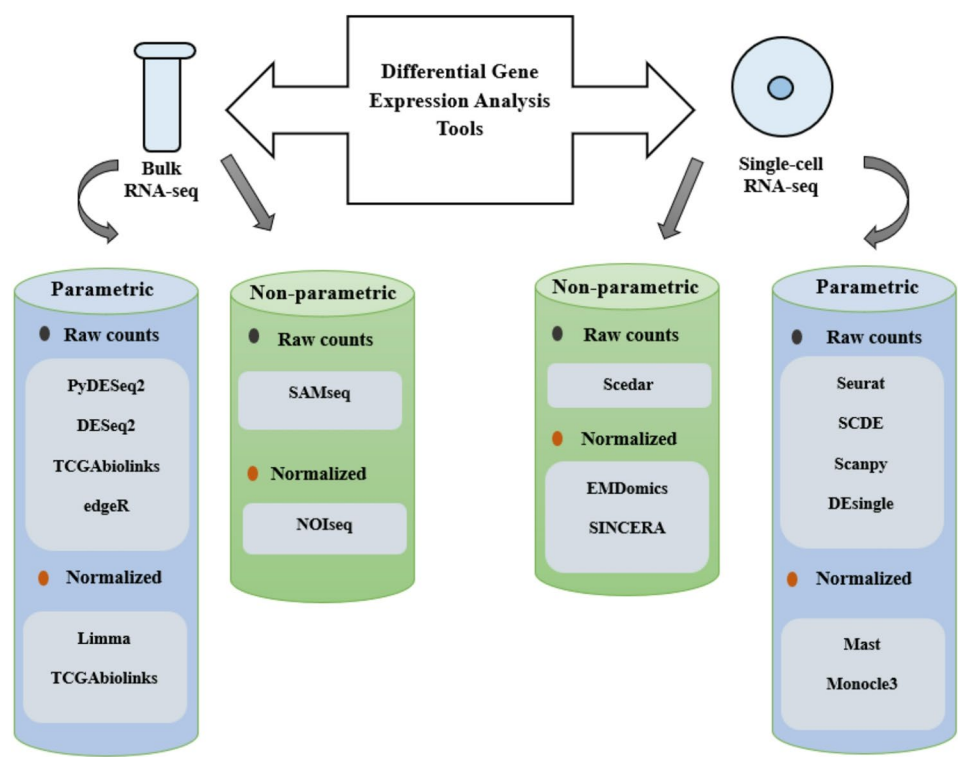


Fig. 2 Categorization of bulk and single cell tools. The packages are grouped based on their DGE methods, which can be either parametric or non-parametric, as well as the type of input data used

Table 1 Comparison of the DEA in the context of bulk, single-cell, direct, and Spatial RNA-seq data interpretation

Tool	Utility	Input	Model	Programming language	Year	Ref.
Limma	Bulk	Norm.	Linear Model	R	2002	[27]
edgeR	Bulk	Counts	NB Model	R	2010	[28]
DESeq2	Bulk	Counts	NB Model	R	2014	[29, 30]
PyDESeq2	Bulk	Counts	NB Model	Python	2022	[37]
NOIseq	Bulk	Norm.	NP	R	2011	[33, 34]
SAMseq	Bulk	Counts	NP	R	2013	[35, 36]
TCGAbiolinks	Bulk	Counts/Norm.	GLM	R	2015	[31]
Monocle2	Single-cell	Norm.	GAM	R	2017	[40]
Mast	Single-cell	Norm.	Generalized Linear Hurdle Model	R	2015	[39]
SCDE	Single-cell	Counts	Mixture of NB & Low-level Poisson	R	2014	[38]
DEsingle	Single-cell	Counts	Zero-Inflated NB	R	2018	[42]
SEURAT	Single-cell	Counts	NB (TCP)	R	2015	[43]
Scanpy	Single-cell	Counts	Graph-based Methods	Python	2018	[47]
Scedar	Single-cell	Counts	KNN	Python	2020	[48]
EMDomics	Single-cell	Norm.	NP	R	2016	[46]
SINCERA	Single-cell	Norm.	NP	R	2015	[45]
Wicoxon	BST	Counts/Norm.	NP	R		
T-test	BST	Norm.	logCPM (TCP)	R		
SpatialDE	ST	Spatial coordinates + counts	Gaussian Process + spatial autocorrelation	Python	2018	[87]
SPARK	ST	Spatial coordinates + counts	GLM with spatial	R	2020	[88]
BayesSpace	ST	Spatial counts	Bayesian hierarchical clustering + spatial smoothing	R	2021	[89]
TAMA	DRS	Counts/transcript	Modular Transcript Annotation and Quantification	Python	2021	[77]
FLAIR	DRS	FASTQ/BAM	Isoform Reconstruction and Splicing-aware Quantification	Python	2020	
TALON	DRS	Counts/Alignments	Gene-/Isoform-level Quantification	Python	2020	[79]
StringTie2	DRS	BAM/Counts	Reference-guided Assembly + Quantification	Python	2019	[80]

Norm. normalized, *NB* negative binomial, *NP* non-parametric, *GLM* generalized linear model, *GAM* generalized additive model, *TCP* two-class parametric, *KNN* K-nearest neighbors, *BST* basic statistical tests, *ST* Spatial transcriptomics, *DRS* Direct RNA sequencing

on identifying genes with spatially variable expression, emphasizing patterns across tissues rather than differential expression between conditions. Because these technologies differ so widely in their capabilities and goals, choosing the right statistical models and computational tools is essential to produce biologically meaningful and reproducible results. The following sections will explore key statistical methods, algorithms, and software designed for DEA across bulk, single-cell, direct RNA, and ST data.

Basic statistical tests for differential expression

Although not specifically designed for RNA-seq count data, several general statistical tests are frequently used as complementary or illustrative approaches in transcriptomic analyses.

Wilcoxon rank sum test (Mann-Whitney test)

The Wilcoxon rank-sum (Mann–Whitney) test is a non-parametric method that assesses whether two independent groups originate from the same underlying distribution by comparing the summed ranks of their observations [30]. While primarily a general statistical test, it is widely adopted in scRNA-seq analyses,

particularly within tools such as Seurat, because of its robustness to sparsity and dropout effects that often challenge count-based parametric models. In this context, the Wilcoxon test provides a simple yet robust approach for detecting differential expression between cell groups, particularly when the data deviate from model assumptions. It can thus serve as a complementary or benchmarking method alongside more sophisticated frameworks [31].

T-test

As a parametric statistical approach, the t-test compares the average values between two groups. It is applicable when specific assumptions are satisfied, including normal data distribution, equal variances, and independence of observations. There are two types of t-tests: independent t-tests for comparing two independent groups and paired t-tests for comparing two dependent groups [32]. Although originally a general-purpose test rather than one tailored for RNA-seq count data, the t-test remains a valuable complementary method in transcriptomic analysis. For example, it is used within frameworks such as limma, where data transformations (e.g., voom) approximate normality and enable parametric modeling. Despite

Table 2 Categorization of techniques employed for identifying DEGs in RNA-sequencing data

	Tool / method	Statistical model / framework	Key strengths & applications	Specific caveats & limitations
General tests	Wilcoxon Rank-Sum Test (Mann-Whitney U)	Non-parametric rank-based test	Distribution-free; robust to outliers. The text notes it controls FDR well in large-scale studies	Ignores magnitude of change; not designed for count data properties (e.g., dispersion, library size); loses power with many tied ranks (common in sparse data)
	Student's t-test	Parametric (assumes normality)	Simple, fast, and robust to mild skewness	Assumptions of normality and equal variance are strongly violated by raw count data and sparse scRNA-seq data
Bulk RNA-seq	DESeq2	Negative Binomial (NB) GLM with shrinkage estimators	Robust dispersion estimation, effective with small replicate counts; handles complex designs; widely adopted	Can be conservative (high false negatives); not designed for the bimodal, zero-inflated nature of scRNA-seq data
	edgeR	NB GLM with empirical Bayes moderation; Exact Test	Powerful for low replicate counts; flexible in experimental design; computationally efficient	Dispersion estimates can be unstable in highly heterogeneous or sparse (e.g., single-cell) data
	limma-voom	Linear model on log-CPM with precision weights	High power; leverages robust linear modeling framework; can model complex designs and batch effects	Assumes the mean-variance trend can be accurately modeled; performance is dependent on the quality of the 'voom' transformation
	PyDESeq2	NB-GLM (Python port of DESeq2)	Native Python integration with Scanpy/AnnData; faster on large datasets	A more recent implementation; may lag behind the R version in features (e.g., certain shrinkage estimators).
Single-cell RNA-seq	MAST	Hurdle Model (Logistic regression for rate + Gaussian for level)	Explicitly models bimodality (dropout rate vs. expression level); can include covariates (e.g., detection rate)	Assumes conditional independence between expression rate and level; ignores the count nature of UMIs (uses log-CPM)
	DEsingle	Zero-Inflated Negative Binomial (ZINB)	Explicitly models dropout events; classifies DE genes into three types (DEs, DEa, DEg) based on expression and dropout rate changes	Computationally intensive; model fitting can be challenging with very low cell counts or extreme sparsity
	SCDE	Bayesian Mixture Model (Poisson for dropouts + NB for expression)	Employs a Bayesian framework to account for technical noise and biological variability	Computationally very intensive; does not scale well to modern large-scale datasets (>10k cells)
	Seurat ('FindMarkers')	Multiple (default: Wilcoxon test). Also supports t-test, MAST, etc.	Part of a comprehensive, widely-used analysis ecosystem; fast and user-friendly	The default Wilcoxon test is not a dedicated scRNA-seq model; it ignores count structure and does not explicitly model zero-inflation
	Scanpy ('rank_genes_groups')	Multiple (default: t-test, Wilcoxon).	Core part of the Python scRNA-seq ecosystem; highly scalable (1 M+ cells); integrates with machine learning tools	Default tests (t-test, Wilcoxon) are general-purpose and do not model single-cell-specific noise structures
	Monocle2 / 3	Generalized Additive Model (GAM)		Not for group-vs-group DE. Identifies genes whose expression changes as a function of pseudotime or across trajectory branches (bifurcations)
Direct RNA-seq	NanoCount + DESeq2/edgeR	Quantification tool + NB-GLM	Leverages robust, mature bulk DE tools (DESeq2, edgeR) after accurate long-read quantification	The final DE analysis is only as good as the initial long-read quantification; does not inherently model long-read-specific error profiles
Spatial transcriptomics	SpatialDE	Gaussian Process Regression	Identifies genes with significant spatial expression patterns (i.e., spatially variable), not just group-vs-group DE	Assumes expression follows a Gaussian distribution after transformation; can be computationally intensive
	SPARK	Generalized Linear Spatial Model (GLSM) with spatial kernels	Robustly handles count data (Poisson/NB); powerful for detecting diverse spatial patterns; accounts for spatial autocorrelation	Computationally demanding, especially for large tissue sections; kernel selection can influence results
	BayesSpace	Bayesian Hierarchical Model	Can enhance the effective resolution of spatial data; identifies spatial domains and clusters; provides probabilistic outputs	Computationally intensive (MCMC-based); results can be sensitive to prior specifications

the high skewness typical of scRNA-seq data, the t-test demonstrates robustness against moderate deviations from normality, occasionally outperforming more complex methods under such conditions [30].

Software packages in R for differential expression analysis of bulk RNA-seq data

Several R-based tools are available for conducting bulk DEA. Below, we outline some of the most widely used packages and provide a comparison of their features in Table 3.

Limma

Limma, a core Bioconductor package, is a widely adopted pipeline for DEA of both microarray and RNA-seq data. For RNA-seq, it is typically combined with the voom transformation, which models the mean-variance relationship of log-counts to produce precision weights for linear modeling. Before applying voom, gene expression counts are typically normalized using the Trimmed Mean of M-values (TMM) method, which is implemented in the edgeR package. The voom transformation converts these normalized counts to the log₂ scale and estimates the mean–variance trend, allowing the application of linear models and empirical Bayes moderation, core strengths of Limma. This approach offers reliable statistical inference and is computationally efficient, making

it especially suitable for studies with small sample sizes. Differential expression significance is assessed using moderated *t*-statistics with Benjamini–Hochberg false discovery rate (FDR) control [33–37]. Limma-voom has become a benchmark method for bulk RNA-seq analyses due to its combination of speed, flexibility in handling complex designs, and effective management of heteroscedasticity through the voom weighting strategy.

EdgeR

The edgeR package provides a statistical framework specifically designed for analyzing replicated count-based gene expression data. It uses the negative binomial (NB) distribution to model RNA-seq counts, allowing for precise estimation of overdispersion caused by biological variability. edgeR supports two main analytical frameworks: the exact test framework, which is suited for simple two-group comparisons (e.g., case vs. control) and is based on a negative-binomial analog of Fisher’s exact test; and the generalized linear model (GLM) framework, which is applicable to complex experimental designs involving multiple factors, interactions, or batch effects.

Both approaches apply empirical Bayes shrinkage to stabilize dispersion estimates across genes, which enhances statistical power while controlling for variability. edgeR also incorporates TMM normalization to correct for library size differences and sequencing depth. FDR control is achieved using the Benjamini–Hochberg procedure [35, 38]. Overall, edgeR provides a robust and flexible framework for performing both simple and complex DEA in bulk RNA-seq experiments.

Table 3 Comparison of DESeq, edgeR, Limma, NOISeq, SAMseq, and PyDESeq2 packages

Tool	Model type	Strengths	Limitations
DESeq2	Parametric (NB)	Robust dispersion modeling; shrinkage estimators; widely used for bulk RNA-seq	Conservative FDR control; not optimized for single-cell
edgeR	Parametric (NB)	Exact tests for small samples; flexible design support	Sensitive to NB assumptions; limited single-cell support
Limma voom	Semi-parametric (voom)	Fast; good for small samples; handles batch effects	Assumes linear mean-variance; bulk-oriented only
NOISeq	Non-parametric	Distribution-free; robust to sequencing depth	Limited to two-group comparisons; less mainstream
SAMseq	Non-parametric	No distributional assumptions; good for sparse data	Slower; underperforms with few DEGs; limited design flexibility
PyDESeq2	Parametric (NB)	Python implementation of DESeq2; reproducible and scriptable	Still maturing; inherits DESeq2’s bulk-focused limitations

NB negative binomial

DESeq2

DESeq2 is a widely used Bioconductor package for DEA based on NB GLMs. It improves on the original DESeq framework (now largely obsolete) by introducing empirical shrinkage estimators for both dispersion and fold-change values. This shrinkage enhances statistical stability, particularly for genes with low read counts or high variability, leading to more reliable detection of DEGs [39, 40]. Normalization is achieved through sample-specific scaling factors that correct for differences in sequencing depth. For hypothesis testing, DESeq2 applies either the Wald test or the likelihood ratio test (LRT), with multiple-testing correction using the Benjamini–Hochberg FDR method [11]. Although optimized for bulk RNA-seq, DESeq2 remains one of the most robust and reproducible parametric tools, balancing conservative FDR control with broad community support and active development.

TCGAbiolinks

TCGAbiolinks is an R/Bioconductor toolkit designed to simplify access, processing, and analysis of data from The

Cancer Genome Atlas (TCGA). Rather than serving as an independent differential expression framework, it primarily acts as a wrapper that integrates well-established Bioconductor tools, such as edgeR and limma-voom, for DEA. The TCGAanalyze_DEA function allows users to perform both simple two-group comparisons and more complex multifactorial designs using these underlying methods. Recent updates have expanded its capabilities to handle multiple covariates and correct for batch effects within TCGA datasets, all while maintaining reproducibility and compatibility with Bioconductor standards [41, 42]. In this way, TCGAbiolinks provides a robust workflow management system that leverages trusted differential expression pipelines within a unified and accessible interface, making it particularly useful for large-scale cancer genomics studies.

Non parametric methods for bulk differential expression analysis

Non-parametric methods offer a robust alternative to parametric frameworks like DESeq2, edgeR, and limma-voom, particularly when the data deviate from model assumptions or when reliable estimation of dispersion and normalization parameters is difficult. Unlike parametric approaches, which rely on specific distributional models (e.g., the NB), non-parametric techniques evaluate differential expression directly from the observed data without assuming an underlying distribution. These methods are especially useful in studies with small sample sizes or heterogeneous datasets, where parametric assumptions may not hold. Rather than replacing traditional parametric tools, non-parametric approaches serve as complementary strategies that enhance analytical robustness under challenging experimental conditions.

NOISeq

NOISeq is an R-based toolkit for RNA-seq analysis encompassing three major components: quality assessment of read counts, filtering and normalization (including batch effect correction), and detection of DEGs [43]. It employs a non-parametric framework that constructs an empirical noise model, providing stable results even when sequencing depth or sample size varies. A key practical strength of NOISeq is its robust performance with very few biological replicates per condition (as few as two or three samples), where parametric models may fail to estimate dispersion reliably. NOISeq generates a reference distribution of noise by pairing fold-change (M) and absolute expression differences (D) among genes within the same condition. This noise distribution, built by pooling all pairwise comparisons among replicates, serves as a baseline to determine whether observed (M, D) values between conditions represent true differential expression rather than random variation [44]. Instead of relying on

a fixed statistical model for read counts, NOISeq determines whether the differences in gene expression are truly significant by comparing them to the actual noise observed in the data. This makes it particularly suited for small-scale or exploratory RNA-seq studies with limited replication, though it is best used as a complementary approach alongside parametric models for validation and cross-verification.

SAMseq

SAMseq is a non-parametric method based on Wilcoxon rank statistics, mainly designed for comparing two groups in RNA-seq experiments. It's particularly robust to outliers and differences in sequencing depth, making it a good choice when normalization or variance modeling might be uncertain. The method uses a resampling approach, repeatedly shuffling the sample labels to adjust for variations in library size and sequencing depth. For each permutation, SAMseq calculates a rank-based statistic, and these results are then combined to produce an overall gene-level score. This process helps control the FDR and ensures reliable results without assuming any particular distribution for gene counts [45, 46]. Although SAMseq and NOISeq share the strength of being resistant to violations of distributional assumptions, they aren't meant to replace parametric methods like DESeq2, edgeR, or limma-voom. Instead, they complement these tools and are especially useful when there are few replicates, when normalization is difficult, or when the data do not fit typical parametric models.

Python-based frameworks for bulk differential expression analysis

In recent years, Python has become one of the most popular languages in bioinformatics, thanks to its flexibility in data processing, smooth integration with machine learning workflows, and an ever-expanding collection of tools for transcriptomic analysis. Within the context of DEA, Python-based frameworks have been developed to complement established R tools and provide accessible workflows for researchers familiar with the Python environment [19].

PyDESeq2

PyDESeq2 is a Python-based implementation of the widely used DESeq2 workflow, developed for analyzing bulk RNA-seq count data. It models raw counts using a NB distribution and fits a GLM for each gene, estimating dispersion parameters and calculating log-fold changes (LFCs) between conditions. By replicating the statistical rigor of DESeq2 within the Python ecosystem, PyDESeq2 allows for seamless integration with other Python-based data science tools. While it currently stands as the main Python solution for bulk RNA-seq DEA, the increasing

adoption of Python in bioinformatics is likely to drive the development of additional, similar frameworks in the near future [47].

Single-cell RNA-seq in R

ScRNA-seq enables gene expression profiling at the resolution of individual cells, allowing researchers to uncover cellular heterogeneity, identify rare cell populations, and investigate dynamic transcriptional states that are often obscured in bulk RNA-seq analyses. Unlike bulk data, scRNA-seq is characterized by high sparsity caused by dropout events, zero inflation, and substantial cell-to-cell variability stemming from both stochastic gene expression and technical noise. These properties necessitate the development of specialized DEA frameworks that can model discrete and continuous components of expression and control for technical artifacts [10]. Several statistical packages have been developed within the R programming environment to address these challenges. For instance, SCDE applies a Bayesian approach to account for dropouts [48], MAST uses a hurdle model to handle bimodal distributions [49], Monocle employs generalized additive models for trajectory-aware DEA [50], DEsingle uses zero-inflated negative binomial modeling [51], and Seurat integrates flexible non-parametric and parametric tests within a unified single-cell analysis workflow [52, 53]. These frameworks collectively form a robust ecosystem for interpreting scRNA-seq datasets and deriving biologically meaningful differential expression results with improved reproducibility and precision.

R packages supporting differential expression analysis in single-cell RNA-seq

SCDE (single-cell differential expression)

SCDE is a computational framework for single-cell RNA-seq analysis that employs a Bayesian probabilistic model to account for both technical noise and intrinsic biological variability. This approach allows for the detection of differential expression patterns and the identification of cell subpopulations with increased noise tolerance. By leveraging individual cell measurements, SCDE estimates gene expression levels across subpopulations and fold changes between them. Compared to traditional RNA-seq methods, SCDE demonstrates enhanced sensitivity, especially for highly expressed genes in specific cell types. Additionally, the package offers error model-based transcriptional similarity metrics and subpopulation analysis [48].

MAST (model-based analysis of single-cell transcriptomics)

MAST is a statistical framework specifically designed for the analysis of scRNA-seq data, employing a hurdle-based GLM optimized for the inherent properties of these datasets. The model decomposes expression into

two components: the probability that a gene is expressed across cells, estimated via logistic regression, and the magnitude of expression among expressing cells, modeled using a Gaussian distribution. Differential expression is subsequently assessed using likelihood ratio testing. By jointly capturing the discrete presence-absence of transcripts and the continuous distribution of expression levels, MAST effectively accommodates the characteristic bimodality of single-cell expression profiles [30, 49].

Monocle2

Monocle is a computational framework for scRNA-seq analysis, developed to detect DEGs across distinct cell populations and along pseudo-temporal trajectories. The method employs a generalized additive modeling strategy, a flexible extension of GLMs in which predictors are expressed as smooth functions to capture non-linear relationships. Monocle2, an enhanced version of the original algorithm, extends its applications to cell type classification, differential expression testing, and trajectory inference. A key feature of Monocle2 is the Census algorithm, which infers relative transcript abundances with greater precision than conventional normalization approaches such as transcripts per million (TPM), thereby improving downstream analytical accuracy [10, 30]. Census counts are easier to model with standard regression techniques, making them a useful tool for DEA. Monocle2 also uses the BEAM (Branch Expression Analysis Modeling) regression model to detect genes that change following fate decisions in development. Monocle evaluates pseudo-time-dependent transcriptional changes by fitting additive models that describe gene expression as a smooth function of pseudo-time, thereby enhancing trajectory-informed differential expression analysis [54, 55].

DEsingle

DEsingle is a statistical package designed for DE analysis in scRNA-seq data, based on a zero-inflated negative binomial (ZINB) modeling framework. The workflow comprises three primary stages: normalization of raw read counts, detection of DEGs, and classification of DEGs into distinct categories. Input data consist of the raw count matrix from scRNA-seq experiments, which is normalized using a modified median approach adapted from DESeq. Parameter estimation for two ZINB populations is then performed through maximum likelihood and constrained maximum likelihood procedures. Differential expression is subsequently tested using likelihood ratio comparisons between the two ZINB populations. Identified DEGs are further subdivided into three categories: DEs, representing genes with significant differences in dropout proportions but not expression levels; DEa, denoting genes with differential expression between groups independent of dropout rates; and

DEg, encompassing genes with significant differences in both dropout proportions and expression abundances. By explicitly modeling zero inflation and distinguishing these expression patterns, DEsingle provides a tailored framework for detecting DEGs in scRNA-seq datasets [30, 51].

SEURAT

Seurat is a comprehensive and widely adopted open-source R framework designed for the preprocessing, analysis, and visualization of scRNA-seq data. Its core functionalities include data quality control, normalization (e.g., LogNormalize or SCTransform), identification of highly variable genes, dimensionality reduction (e.g., PCA, UMAP, or t-SNE), graph-based clustering, DEG testing (using methods such as Wilcoxon rank-sum, likelihood ratio, or ROC tests), and cell-type annotation based on marker gene expression [52, 56, 57]. Seurat also provides advanced methods for integrating multiple datasets or batches using anchor-based approaches, allowing alignment of shared cell populations across different conditions, technologies, or species [57]. Extending beyond traditional single-modality analyses, Seurat has developed into a comprehensive multimodal framework that enables the integrated analysis of diverse single-cell data types, such as RNA, protein, and chromatin accessibility, as well as emerging spatial transcriptomics datasets. In its spatial analysis modules, Seurat integrates gene expression matrices with spatial coordinates and histological images (e.g., from 10x Visium or Xenium platforms) to enable visualization and exploration of gene expression in spatial context [56]. Importantly, while Seurat supports spatial transcriptomic data visualization and analysis, its primary purpose remains the comprehensive processing, integration, and interpretation of scRNA-seq and multimodal single-cell datasets rather than the direct spatial reconstruction of cellular positions from gene expression profiles.

Single-cell gene expression profiling in python

Python has become a powerful platform for scRNA-seq data analysis, valued for its versatility, scalability, and smooth integration with machine learning frameworks. Unlike bulk or microarray expression experiments, scRNA-seq can involve thousands to millions of individual cells, demanding efficient handling of sparse, high-dimensional data and robust computational tools. At the core of many Python-based toolkits lies the AnnData data structure, which provides a flexible framework for storing and managing large sparse matrices along with associated metadata. The widely used Scanpy package provides a comprehensive pipeline for scRNA-seq analysis, encompassing quality control, normalization, dimensionality reduction, clustering, visualization (e.g., UMAP,

t-SNE), and DEG testing [58]. Advanced frameworks such as Scvi-tools apply deep generative models for batch correction, data integration, and probabilistic DEA, harnessing variational autoencoders to model complex biological variation [59]. The Scedar package supports exploratory data analysis, including clustering and dropout imputation, optimized for high-throughput single-cell datasets [60]. Taken together, these Python resources establish a comprehensive ecosystem for scalable, reproducible, and interpretable analysis of single-cell transcriptomic data.

Python implementations for differential expression analysis in single-cell datasets

Scanpy

Scanpy is a comprehensive Python-based toolkit widely adopted for single-cell transcriptomic analysis, featuring robust capabilities for DEG testing. Its core functionality for DEA is implemented through the `tl.rank_genes_groups` module, which enables statistical testing between clusters or predefined groups of cells using methods such as the t-test, Wilcoxon rank-sum test, and logistic regression. These functions operate on both raw and normalized expression values and scale efficiently to large datasets through the optimized AnnData data structure. While Scanpy provides extensive tools for clustering, visualization, and trajectory inference, its primary role in DEA is to identify genes that distinguish specific cell populations or states after clustering or annotation has been completed. In addition, Scanpy can interface with `diffxpy`, a companion library that applies GLMs to gene expression data, offering greater statistical flexibility for complex experimental designs. Moreover, Scanpy integrates smoothly with the broader Python data science and machine learning ecosystem, including frameworks such as TensorFlow and PyTorch, supporting scalable and extensible workflows for large-scale single-cell data analysis. Overall, Scanpy provides an efficient and versatile framework for identifying DEGs across large single-cell cohorts while maintaining compatibility with advanced statistical and deep learning tools [58, 61].

Scedar

Scedar is primarily designed for exploratory analysis, clustering, and visualization of single-cell RNA-seq datasets rather than for direct differential expression testing. It offers an interactive computational framework that incorporates KNN-based operations, clustering algorithms, and visualization tools, allowing users to investigate cellular heterogeneity and identify genes that distinguish between clusters. Although Scedar does not include dedicated statistical testing modules for DEA, its visualization and clustering capabilities can reveal biologically meaningful cell groups and gene expression patterns, which can subsequently be subjected to differential

analysis using frameworks such as Scanpy or diffxpy. Thus, Scedar serves as an exploratory analysis platform that complements dedicated DEA tools by helping to define cell clusters and relationships for subsequent differential testing [60].

Diffxpy

Diffxpy is a Python-based framework designed specifically for DEA in scRNA-seq data. It provides a versatile set of statistical models, such as Wald tests, likelihood-ratio tests (LRTs), and models that handle continuous covariates, making it applicable to both discrete and continuous experimental conditions. Unlike simpler non-parametric methods, diffxpy uses GLMs with customizable covariates and design matrices, allowing researchers to account for confounding factors and complex experimental setups. The framework integrates seamlessly with Scanpy and its AnnData structure, ensuring smooth compatibility with common single-cell analysis workflows. Emphasizing both scalability and interpretability, diffxpy delivers efficient computation and detailed model outputs, making it a practical tool for large-scale single-cell studies (Theislab, diffxpy) [61].

scvi-tools

scvi-tools extends DEA into the probabilistic modeling domain using deep generative models powered by deep learning. Built on variational autoencoders, it performs Bayesian inference to estimate latent representations of gene expression while simultaneously correcting for batch effects and other technical confounders. This approach helps identify genes that are expressed differently by using a probabilistic framework that naturally accounts for uncertainty in the parameter estimates. scvi-tools is particularly powerful for multi-condition and multi-batch datasets, providing statistically rigorous and scalable inference across large cohorts of cells. Its modular design also allows it to integrate with other omics data types, making it a flexible and powerful tool for modern single-cell analysis [59].

Non parametric methods for differential gene activity assessment across single-cell

SINCERA: single cell RNA-seq profiling analysis

SINCERA is one of the early computational frameworks developed for downstream single-cell transcriptomic analysis, encompassing preprocessing, normalization, cell type classification, DEA, gene signature inference, and transcriptional regulator identification. For DEG testing, SINCERA applies statistical hypothesis testing to compute gene-level p-values between two groups using either a one-tailed Welch's t-test, assuming independent normal distributions, or the non-parametric Wilcoxon rank-sum test, which is more suitable for smaller sample

sizes. FDRs are managed using the Benjamini–Hochberg correction. While SINCERA played an important role in establishing foundational methodologies for single-cell RNA-seq analysis, it is now considered primarily of historical and methodological interest. Modern workflows like Seurat, MAST, and Scanpy include advanced normalization methods, model-based variance control, and improved scalability for handling large datasets. Consequently, SINCERA is less commonly used in contemporary single-cell pipelines but remains valuable as one of the first integrated frameworks for comprehensive scRNA-seq analysis [10, 62].

EMDomics

EMDomics is another early non-parametric approach for DEA, which employs the Earth Mover's Distance (EMD) to quantify divergence between gene expression distributions. For each gene, an EMD score is calculated, and the FDR is estimated to determine statistical significance. Unlike traditional DEA methods that compare average expression levels, EMDomics captures distribution-level differences, offering sensitivity to subtle shifts in expression heterogeneity across cells. Although methodologically innovative, EMDomics has seen limited adoption in recent scRNA-seq studies due to scalability constraints and the availability of more flexible, model-based frameworks such as MAST, Seurat, and Scanpy, which integrate distributional modeling and variance correction directly into their pipelines. Nevertheless, EMDomics represents an important conceptual milestone, introducing distributional comparisons as a meaningful statistical paradigm in transcriptomic data analysis [10, 63].

Comparison of packages for single-cell differential expression analysis

To evaluate the performance of various frameworks for single-cell DEA, we compared representative tools based on reported benchmarking studies [64, 65]. Earlier analyses indicated that Monocle2 and EMDomics identified large numbers of true DEGs but also produced relatively high false-positive rates, leading to moderate overall accuracy (0.824 and 0.91, respectively) [64]. However, both methods have since been superseded by more recent and robust implementations (e.g., Monocle3) or replaced by modern alternatives such as Seurat, MAST, and Scanpy, which better handle sparsity and zero inflation in single-cell data [52, 65]. Tools with higher precision, such as MAST, SCDE, edgeR, and SINCERA, tended to produce fewer false positives but also detected fewer true positives [1]. Although DESeq2 and edgeR, originally developed for bulk RNA-seq analysis, have been adapted for single-cell datasets, they do not explicitly model zero inflation or dropout effects, limiting their accuracy in this context [65]. Among the evaluated methods, DEsingle

Table 4 Summary comparison of single-cell packages for diverse aspects of scRNA-seq analysis

Category	Tool	Language	Primary function	Relevance to DEG analysis
Annotation / identity	SingleR	R	Automated annotation	Supports DEG interpretation
	SCENIC	R	Regulatory network inference	Links regulators to DEGs
Trajectory / pseudotime	Slingshot	R	Lineage and pseudotime inference	Identifies dynamic DEGs
Regulatory / signaling	SingleCell-SignalR	R	Cell-cell communication inference	Context for DEGs
Multi-omics / large-scale	MultiVelo	Python	Chromatin-integrated velocity	Regulatory insight on DEGs
	SnapA-TAC2	Python	Multi-omics integration	Cross-layer DEG patterns
	Mowgli	Python	Joint transcriptome–epigenome analysis	Integrative DEG interpretation
	Cumulus	Python	Cloud-based large-scale analysis	Scalable DEG pipelines

demonstrated the highest overall accuracy and F1 score, achieving a favorable balance between sensitivity and precision [51].

Single-cell packages for diverse systematic approaches to scRNA-seq data interrogation

ScRNA-seq analysis involves multiple computational tasks beyond standard DEG detection such as normalization, cell-type identification, trajectory inference, and integrative multi-omics analysis.

To improve clarity and emphasize DEA, the next section organizes key single-cell analysis packages into thematic categories based on their main analytical goals. Table 4 summarizes their core functionalities, programming language, and relevance to differential expression analysis.

Annotation and cell-type identification

SingleR is an automated framework for annotating scRNA-seq data, which assigns cell identities in a query dataset by leveraging reference collections with established labels. Instead of requiring repeated manual inspection of clusters and redefinition of marker genes for each experiment, the interpretive effort is consolidated into the construction of the reference, and the embedded biological knowledge is subsequently transferred to new datasets in an automated manner. When

integrated with Seurat, a widely adopted platform for scRNA-seq data processing and analysis, *SingleR* provides a synergistic strategy that streamlines cell-type identification and enhances downstream biological interpretation. By capitalizing on reference transcriptomic resources, the method enables robust and reproducible recognition of cellular phenotypes in scRNA-seq data, thereby offering a valuable asset for researchers engaged in high-resolution cellular profiling [66].

SCENIC (Single-Cell Regulatory Network Inference and Clustering) provides a framework in R for combining single-cell transcriptomic data with gene regulatory network analysis. It allows for elucidating transcriptional regulators and regulatory networks driving cell identity and differentiation. The SCENIC workflow involves three main steps. First, co-expression modules are identified using GENIE3, which may contain false positives and indirect targets. Next, the co-expression modules are refined by analyzing them integrating cis-regulatory motif identification via RcisTarget to detect candidate direct transcription factor targets. Finally, regulator activity in each cell is quantified using AUCell, enabling the identification of cells exhibiting significantly elevated subnetwork activity. The derived binary matrix enables downstream applications, including clustering, to resolve cell types and states guided by shared regulatory network activity. This approach is robust against drop-outs since the regulon is scored as a whole rather than individual genes or transcription factors [67].

Trajectory and pseudotime inference

Slingshot is an R platform designed to delineate cellular developmental lineages and infer pseudo-temporal progression from scRNA-seq datasets. It addresses the challenges of noisy, multidimensional single-cell expression profiles by providing a flexible approach that allows for the discovery of cell fate landscapes sculpted by biologically principled structural assumptions, and seamless integration with widely adopted analytical infrastructures. The methodology involves two main stages: the identification of lineages using a topology-preserving tree constructed from clustered cellular states and the inference of pseudo-times using simultaneous principal curves. *Slingshot*’s flexibility in accommodating various upstream data handling strategies, encompassing normalization schemes, embedding transformations, and cell-type resolution methods,, makes it a valuable tool for researchers within the landscape of cellular-resolution transcriptomics. Additionally, it provides the option accommodates localized supervision by allowing users to specify initial and terminal cellular states, thereby improving the resolution and biological plausibility of the inferred global lineage topology. This approach is operationalized within the open-source R package

Slingshot, which is distributed via the Bioconductor platform, ensuring accessibility and seamless integration into established single-cell analysis workflows [68].

Regulatory network and signaling inference

SCENIC, beyond its clustering role, refines co-expression networks through motif enrichment (RcisTarget) and activity scoring (AUCell), enabling detection of key transcriptional regulators underlying DEGs [67].

SingleCellSignalR is an R-based computational framework designed to reconstruct intercellular communication networks from single-cell transcriptomic data. It incorporates a curated repository of experimentally validated ligand–receptor interactions (LRdb) and introduces a regularized product scoring scheme to accommodate variability in sequencing depth across datasets. Starting from raw read count matrices, the package provides an integrated workflow that encompasses data normalization, unsupervised clustering, and automated cell-type annotation, thereby streamlining the inference of cell–cell signaling landscapes. *SingleCellSignalR* allows for the inference of LR interactions between cell populations and provides various analysis results in tabular and graphical formats. It also enables the visualization of intercellular networks and provides functionality for receptor downstream signaling exploration via integration of Reactome and KEGG pathways, with an open architecture that allows interaction with other packages and adaptability to single-cell proteomics technologies. Overall, *SingleCellSignalR* provides an important approach for resolving cellular communication maps and nominating ligand–receptor interactions with sufficient confidence to warrant further validation [69].

Multi omics and large-scale frameworks

MultiVelo is a computational framework that integrates chromatin accessibility data into the modeling of single-cell gene expression dynamics. *MultiVelo* extends the RNA velocity framework by incorporating epigenomic data, providing more accurate predictions of lineage outcomes as opposed to RNA-only velocity-driven models. The model uses a probabilistic latent variable approach to infer the chronological interplay between epigenomic and transcriptomic dynamics, revealing distinct classes of genes based on the timing of chromatin closure relative to transcription cessation. *MultiVelo* is demonstrated to be effective in various tissues and cell types, offering insights into gene regulation and cell fate determination. The study's findings suggest that incorporating chromatin accessibility into gene expression models markedly enhances velocity inference, especially within early stem cells experiencing rapid epigenomic remodeling [70].

SnapATAC2 is a versatile and scalable software package to interrogate single-cell omics landscapes. The

framework enables a fast and efficient tool for researchers, with key features including a preprocessing module for handling large-scale datasets, leveraging a matrix-free spectral embedding paradigm, the algorithm achieves high-fidelity dimensionality reduction across heterogeneous single-cell modalities and enables cohesive integration of multi-omics layers into a biologically meaningful latent manifold, and a distributed and parallelized approach for overcoming memory constraints. This software marks a substantial leap in single-cell data analytics, delivering a robust, scalable, and user-oriented platform for decoding gene regulatory architectures through integrative analysis of single-cell genomic, transcriptomic, and epigenomic landscapes [71].

Mowgli is a computational framework developed for the joint analysis of paired single-cell multi-omics data, with the aim of improving both clustering accuracy and biological insight. The approach integrates Nonnegative Matrix Factorization (NMF) with Optimal Transport (OT) to derive cell type–resolved factors from TEA-seq measurements, thereby surpassing the interpretability achieved by current state-of-the-art strategies. Comprehensive benchmarking against widely adopted platforms, including MOFA + and Seurat v4, demonstrates that *Mowgli* attains comparable or superior performance while offering enhanced mechanistic clarity. The method is distributed as a Python package, fully compatible with the scverse ecosystem, and designed for straightforward installation. To facilitate transparency and reproducibility, all scripts and resources used to generate the analyses and figures are publicly available via GitHub. Released under a CC-BY-NC-ND 4.0 International license, *Mowgli* has been successfully applied across multiple datasets, underscoring its utility in advancing integrative single-cell multi-omics research [72].

Cumulus is a cloud-based platform engineered to enable the analysis of large-scale single-cell and single-nucleus RNA-seq datasets. By integrating cloud computing with algorithmic innovations, it achieves high scalability, cost efficiency, and streamlined usability, while offering comprehensive functionality. The framework encompasses three principal components: a cloud-optimized analysis pipeline, the Pegasus Python package for downstream computational tasks, and Cirrocumulus, an interactive visualization environment. It is compatible with multiple data modalities and supports end-to-end workflows ranging from raw read processing to advanced biological interpretation. Benchmarking indicates that *Cumulus* delivers greater efficiency than conventional frameworks without compromising, often even improving, the quality of results, thereby rendering it highly suitable for population-scale studies in single-cell genomics. The analytical workflow proceeds through three core stages: (1) sequence read processing, (2) construction of

gene-count matrices, and (3) biological data interrogation. Its design leverages the synergy of cloud resources, optimized algorithms, and efficient implementation strategies to overcome the computational challenges posed by large-scale sc/snRNA-seq experiments. While cloud-enabled approaches for scRNA-seq analysis have been explored previously, Cumulus represents the most comprehensive and fully integrated solution to date [73].

Direct RNA sequencing (DRS)

Direct RNA Sequencing (DRS) is an advanced transcriptomic approach that enables the sequencing of nearly full-length RNA molecules without the need for cDNA conversion or amplification [21]. Enabled by Oxford Nanopore Technologies (ONT), DRS allows direct profiling of native RNA molecules, thus preserving critical epitranscriptomic information and facilitating accurate isoform identification and alternative splicing analysis [22]. In nanopore sequencing, RNA strands are passed through a nanoscale biological pore embedded in a membrane, and changes in ionic current are decoded by a basecaller to determine the nucleotide sequence [21]. In contrast to short-read sequencing technologies that rely on reverse transcription and PCR, DRS avoids amplification bias and retains native RNA modifications, making it highly advantageous for transcriptome analysis at per-read, per-site resolution. Recent developments such as the Dorado basecaller [74] have significantly improved base-level accuracy (~ 98%) and enabled *de novo* detection of common RNA modifications including pseudouridine, N6-methyladenosine, 5-methylcytosine, and inosine [75, 76].

While DRS excels at transcript-level resolution, its application to differential gene expression analysis presents unique computational challenges due to long-read characteristics and inherent sequencing noise. These challenges mainly arise in upstream processes such as read assignment, isoform disambiguation, and error correction, which directly influence the accuracy of downstream quantification. Specialized tools have been developed or adapted to address these complexities. NanoCount is a Python-based quantification tool optimized for ONT reads, offering robust transcript-level abundance estimates for downstream analysis [77]. TAMA (Transcriptome Annotation by Modular Algorithms) supports transcript quantification and DEA pipelines tailored to long-read transcriptomes [78]. In addition, tools such as FLAIR [79], TALON [80], StringTie2 (long-read version) [81], and IsoQuant [82] are widely used for isoform-level reconstruction and quantification, providing essential input for subsequent DEA workflows. Widely used statistical packages such as edgeR and DESeq2 have been successfully applied to count matrices derived from these quantification tools;

however, they are not DRS-specific and rely on the accuracy of upstream transcript reconstruction.

Beyond tool development, ONT provides workflows through its EPI2ME platform [83] that integrate commonly used methods such as DESeq2 and edgeR, although several open-source community pipelines and benchmarking studies now offer comparable alternatives [84]. Comparative analyses between DRS and short-read RNA-seq have demonstrated that while DRS improves isoform-level resolution and enables the detection of RNA modifications, it may show greater variability in quantification. Therefore, benchmarking against established short-read datasets and the use of negative controls remain essential to ensure reproducibility and biological interpretability. As DRS adoption grows, continued refinement of computational pipelines and benchmarking efforts will be critical for achieving robust and reproducible differential expression analysis from long-read transcriptomic data.

Spatial transcriptomics

Spatial transcriptomics is an innovative approach that combines gene expression profiling with spatial information, allowing researchers to map transcriptomic data directly onto tissue architecture [85]. Unlike bulk and single-cell RNA sequencing, which require tissue dissociation and thereby lose spatial context, spatial transcriptomics preserves tissue morphology and the native positioning of cells within their microenvironment [26]. This enables the simultaneous investigation of gene expression patterns and tissue structure, offering deep insights into cell-cell interactions, tissue heterogeneity, and the functional organization of complex tissues. Such spatial resolution is particularly valuable in cancer research [86], neuroscience, and developmental biology, where spatially regulated gene expression plays a critical role in biological function and disease progression [87].

For differential gene expression analysis, spatial transcriptomics allows comparison of gene expression across defined regions or tissue zones, accounting for spatial dependencies that are absent in dissociated single-cell or bulk datasets. Several computational tools have been developed in both R and Python to perform DGE analysis using spatially resolved transcriptomic data. SpatialDE (Python) employs statistical models to identify spatially variable genes by capturing spatial autocorrelation in gene expression patterns [88]. SPARK (R) offers a robust framework for detecting spatially DEGs using kernel-based methods [89]. BayesSpace (R) leverages Bayesian statistical modeling to detect gene expression variability across spatial domains with high resolution [90]. The Seurat package (R), widely known for scRNA-seq analysis, includes a spatial module that supports clustering and DEA in spatial transcriptomic contexts [52], while

Scanpy (Python) offers similar functionality through its spatial module, enabling clustering and DEA analysis within spatially annotated datasets [58]. Giotto (R) provides a comprehensive and interactive platform that incorporates spatial information to refine DEG detection and visualize spatially distinct expression patterns [91].

Recent benchmarking studies have compared these methods, showing that performance can vary depending on spatial resolution, tissue complexity, and normalization strategies [92]. For instance, SPARK and BayesSpace have demonstrated improved sensitivity for detecting spatially varying genes in structured tissues, while SpatialDE performs consistently across a range of experimental designs [88]. However, challenges remain due to technical noise, batch variability, and the computational complexity of integrating histological imaging with gene expression data [93]. Additionally, resolution limitations and spot-level averaging in technologies such as 10x Visium can mask cell-specific expression differences, complicating DEG interpretation. The high cost of spatial transcriptomics experiments and the need for specialized imaging platforms further constrain accessibility. As the field evolves, continuous benchmarking using standardized datasets and open-access resources will be essential to ensure reproducibility, comparability, and reliability of DEA results across platforms.

Challenges in differential expression analysis across RNA-seq modalities

DEA from RNA-seq data presents several technical and computational challenges. One major issue is normalization, as differences in sequencing depth, library composition, or RNA composition can bias comparisons across samples. Tools such as DESeq2 [40] and edgeR [38] implement normalization strategies like median-of-ratios or trimmed mean of M-values to mitigate these effects. Batch effects and other technical confounders add further complexity to the analysis, as uncontrolled variability can obscure or imitate genuine biological signals. Statistical frameworks such as limma [34] and ComBat [94] are commonly applied to correct for these issues. Additionally, lowly expressed genes and small sample sizes increase statistical uncertainty, heightening the risk of both false positives and false negatives. In addition, error rates in read quantification, stemming from alignment ambiguities or biases introduced during library preparation and cDNA synthesis, can propagate into DGE results [95, 96]. Addressing these challenges requires the careful selection of normalization methods, statistical models, and computational tools tailored to the specific RNA-seq modality, ensuring robust and reproducible identification of DEGs.

Challenges of single cell RNA-Seq for DEA

Analyzing scRNA-seq data to find DEGs comes with quite a few challenges. One big problem is elevated non-biological signal variance in single-cell datasets, caused by things like amplification biases and batch effects. This noise can mask the real biological signals, making it hard to pinpoint true DEGs accurately [97]. Another challenge is the sheer scale of the data, each cell can have expression measurements for an extensive repertoire of expressed genes across the genome, which makes the analysis very complex and increases the chance of false positives [98]. The data is also very sparse, meaning many genes show zero expression in a lot of cells. This scarcity makes it tough to get reliable gene expression estimates and contributes to more false discoveries [99]. Batch effects add another layer of difficulty since experiments are often run in multiple batches, bringing in systematic biases that are tricky to correct fully [100]. Moreover, the intrinsic cell-to-cell variability captured by scRNA-seq reflects genuine biological heterogeneity but also complicates statistical confidence in identifying DEGs, as this variability can mask consistent expression differences [101]. Preprocessing is a critical step for successful DEA. It usually involves filtering out low-quality or outlier cells, if you keep these, they can skew the results. Another practical issue is that analyzing large scRNA-seq datasets takes a lot of computing power and time. Many current methods rely on complex statistical models applied to each gene with iterative algorithms, which can take hours, even for smaller datasets, and much longer for bigger ones. Lastly, false discoveries are a serious concern in single-cell DEA. If methods don't properly account for the cell-to-cell variation, there's a risk of identifying genes as differentially expressed when they're not, leading to more false positives compared to bulk RNA-seq studies. Recent research has pointed out this problem, highlighting the importance of careful computational strategies to reduce false discoveries and make DEA more reliable in single-cell data [20].

Challenges of direct RNA sequencing (DRS)

Direct RNA sequencing offers substantial advantages over traditional short-read RNA-seq by eliminating the need for cDNA synthesis and amplification, thereby preserving full-length transcript structures and enabling the detection of native RNA modifications [75]. However, despite these benefits, DRS introduces a distinct set of challenges that complicate its application in differential gene expression analysis [102–104].

From an experimental standpoint, DRS is still limited by error rates compared to short-read sequencing, which can affect both basecalling and modification detection accuracy. While recent improvements in nanopore chemistry and the Dorado basecaller have enhanced signal

resolution, distinguishing true RNA modifications from sequencing noise remains difficult, particularly in low-coverage regions [105]. The high cost of nanopore flow cells, consumables, and instrument maintenance further constrains large-scale transcriptomic applications, especially when high read depth is required for robust quantification [102–104]. Moreover, DRS throughput remains lower than that of conventional short-read sequencing, limiting its scalability for population-level studies.

From a computational perspective, the long-read and error-prone nature of DRS data presents challenges in alignment, transcript quantification, and isoform disambiguation. Standard DEA tools originally developed for short-read sequencing often fail to provide consistent results when applied to long-read data, leading to inaccuracies in expression quantification. Although specialized tools such as NanoCount [77] and TAMA [78] have been developed to support transcript quantification and DEA from DRS data, further validation is still needed to reach the reliability of short-read methods like DESeq2 and edgeR. Additional long-read tools such as FLAIR, TALON, StringTie2, and IsoQuant have improved isoform-level quantification but remain computationally intensive and require substantial memory and runtime resources. Many widely used RNA-seq pipelines are not yet optimized for long-read data, underscoring the need for standardized workflows that integrate long-read-specific preprocessing, alignment, and quantification. Together, these challenges highlight the necessity for continued tool development, benchmarking, and community-wide standardization to fully realize the potential of DRS in transcriptomic research.

Challenges of spatial transcriptomics

Spatial transcriptomics offers a significant advancement in transcriptomic profiling by preserving spatial context within tissues, providing insights into cell-type localization and interactions that are lost in bulk and single-cell RNA-seq approaches [106–109]. However, this technology also presents several technical and computational challenges that complicate differential gene expression analysis.

From a technological perspective, one major limitation is the resolution variability among platforms; for example, some technologies such as 10x Visium capture transcriptomes at a multi-cellular resolution rather than achieving true single-cell granularity, potentially masking finer spatial heterogeneity [85]. Emerging high-resolution platforms such as STOmics (Stereo-seq) [110, 111] and CosMx Spatial Molecular Imaging (SMI) [112] now allow near single-cell or even subcellular resolution, addressing some of these limitations but introducing new challenges in data volume, imaging noise, and cost. Additionally, tissue sectioning procedures can introduce sampling biases

that reduce reproducibility and influence DEG signals across spatial coordinates. Batch effects and artifacts from tissue preparation or imaging add another layer of variability, making it difficult to compare samples or integrate results across experiments.

From a computational standpoint, performing accurate DEA in spatial transcriptomics requires statistical frameworks that account for spatial autocorrelation, dependencies between neighboring spatial units, that traditional bulk RNA-seq methods do not model. Tools like SPARK [89], SpatialDE [88], and BayesSpace [90] provide spatially aware DEA, but their outputs still require careful biological validation to ensure accurate interpretation. Another challenge lies in the presence of batch effects and technical variability between tissue sections, which can introduce systematic biases and hinder cross-experiment comparisons. Methods like Seurat's integration framework and Scanpy's spatial modules aim to correct for these issues, but a universally accepted standard for batch correction in spatial data is lacking. Many DEA frameworks for spatial data still rely on scRNA-seq-based normalization and modeling assumptions, which may bias results by ignoring spatial dependency structures. Recent benchmarking efforts comparing these methods have shown that performance varies significantly depending on tissue type, platform resolution, and preprocessing strategy, emphasizing the need for standardized evaluation metrics and datasets [106–109].

The high cost and complexity of spatial transcriptomics further limit its accessibility. These experiments require specialized imaging platforms, sequencing protocols, and computational workflows to integrate spatial transcript data with histological context, raising both financial and technical barriers [106–109]. Moreover, analyzing spatial transcriptomics data demands the integration of multiple data types, such as histological images, spatial gene expression matrices, and potentially matched scRNA-seq data. This makes the computational pipelines more sophisticated and resource-intensive. Developing unified frameworks that combine DEA, spatial patterning, and image-derived features remains an active area of research. Continued benchmarking and collaboration between computational and experimental groups will be essential to improve reproducibility, comparability, and biological interpretability in spatial DEA studies.

Conclusion

RNA-seq technologies have revolutionized DEA by providing deeper insights into gene regulation across diverse biological systems. Bulk RNA-seq remains a reliable and cost-effective approach for detecting DEGs at the population level, offering high reproducibility and well-established analytical workflows. In contrast, scRNA-seq enables the study of cellular heterogeneity and the

identification of rare cell populations, although it continues to face challenges such as dropout events, data sparsity, and variability in gene expression. Emerging modalities like DRS facilitate the direct detection of RNA isoforms and modifications, while spatial transcriptomics preserves tissue architecture, revealing spatial patterns of gene expression within their native context.

Despite these advancements, several challenges persist. The integration of multimodal datasets, refinement of statistical models to better handle sparsity and technical noise, and the development of robust benchmarking frameworks are essential for improving the accuracy of DEA. Continued progress in computational methods and sequencing technologies will further enhance the resolution, precision, and biological relevance of RNA-seq studies, leading to a more comprehensive understanding of gene regulation and its impact on health and disease.

Acknowledgements

This work was supported in part by Shahrekord University [1403].

Author contributions

MKD was responsible for formal analysis, experimental investigation, methodological design, resource curation, software implementation, data validation, and visualization. MK also drafted the initial manuscript and participated in its critical revision and editorial refinement. NGE performed formal analysis, contributed to experimental investigation, methodological development, resource management, software execution, and validation procedures. NG likewise authored the original draft and engaged in manuscript review and editing. MEB led the conceptual framework of the study and contributed to formal analysis, funding acquisition, experimental design, investigation, and methodological oversight. ME also managed project administration and provided supervisory guidance. In addition, ME authored the initial draft and contributed substantively to the manuscript's review and editorial process. All authors read and approved the final manuscript.

Funding

Shahrekord University.

Data availability

No datasets were generated or analysed during the current study.

Declarations

Ethics approval and consent to participate

Not applicable.

Consent for publication

Not applicable.

Competing interests

The authors declare no competing interests.

Author details

¹Department of Genetics, Faculty of Basic Sciences, Shahrekord University, Shahrekord, Iran

²Department of Bioengineering, Northeastern University, Boston, MA, USA

References

1. Conesa A, et al. A survey of best practices for RNA-seq data analysis. *Genome Biol.* 2016;17(1):13.
2. Mortazavi A, et al. Mapping and quantifying mammalian transcriptomes by RNA-Seq. *Nat Methods.* 2008;5(7):621–8.
3. Corchete LA, et al. Systematic comparison and assessment of RNA-seq procedures for gene expression quantitative analysis. *Sci Rep.* 2020;10(1):19737.
4. Hong M, et al. RNA sequencing: new technologies and applications in cancer research. *J Hematol Oncol.* 2020;13(1):166.
5. Stupnikov A, et al. Robustness of differential gene expression analysis of RNA-seq. *Comput Struct Biotechnol J.* 2021;19:3470–81.
6. Yang IS, Kim S. Analysis of whole transcriptome sequencing data: workflow and software. *Genomics Inf.* 2015;13(4):119.
7. Jaryani F, Adhikar B, Beheshti S et al. Sixth annual BCM hackathon on structural variation and pangenomics [version 1; peer review: awaiting peer review]. *F1000Research.* 2025;14:1231. <https://doi.org/10.12688/f1000research.170665.1>.
8. Li D. Statistical methods for RNA sequencing data analysis. Exon; 2019. pp. 85–99.
9. Costa-Silva J, Domingues D, Lopes FM. RNA-Seq differential expression analysis: an extended review and a software tool. *PLoS ONE.* 2017;12(12):e0190152.
10. Wang T, et al. Comparative analysis of differential gene expression analysis tools for single-cell RNA sequencing data. *BMC Bioinformatics.* 2019;20(1):40.
11. Alessandri L, Arigoni M, Calogero R. Differential expression analysis in single-cell transcriptomics. In: *Single cell methods: sequencing and proteomics.* Springer; 2019. pp. 425–432.
12. Kolodziejczyk AA, et al. The technology and biology of single-cell RNA sequencing. *Mol Cell.* 2015;58(4):610–20.
13. Van den Berge K, et al. Observation weights unlock bulk RNA-seq tools for zero inflation and single-cell applications. *Genome Biol.* 2018;19(1):24.
14. Chen G, Ning B, Shi T. Single-cell RNA-seq technologies and related computational data analysis. *Front Genet.* 2019;10:317.
15. Miao Z, Zhang X. Differential expression analyses for single-cell RNA-Seq: old questions on new data. *Quant Biology.* 2016;4(4):243–60.
16. Sengupta D et al. Fast, scalable and accurate differential expression analysis for single cells. *BioRxiv.* 2016; 49734.
17. Das S et al. A comprehensive survey of statistical approaches for differential expression analysis in single-cell RNA sequencing studies. *Genes.* 2021;12(12):1947.
18. Jaakkola MK, et al. Comparison of methods to detect differentially expressed genes between single-cell populations. *Brief Bioinform.* 2017;18(5):735–43.
19. Bystrykh L. Python for gene expression. *F1000Research.* 2022;10:870.
20. Das S, Rai A, Rai SN. Differential expression analysis of single-cell RNA-Seq data: current statistical approaches and outstanding challenges. *Entropy.* 2022;24(7):995.
21. Garalde DR, et al. Highly parallel direct RNA sequencing on an array of nanopores. *Nat Methods.* 2018;15(3):201–6.
22. Workman RE, et al. Nanopore native RNA sequencing of a human Poly (A) transcriptome. *Nat Methods.* 2019;16(12):1297–305.
23. Stephenson W, et al. Direct detection of RNA modifications and structure using single-molecule nanopore sequencing. *Cell Genom.* 2022; 2(2):100097. Stephenson W et al. Direct detection of RNA modifications and structure using single-molecule nanopore sequencing. *Cell Genomics.* 2022;2(2). <https://doi.org/10.1016/j.xgen.2022.100097>
24. Leger A, et al. RNA modifications detection by comparative nanopore direct RNA sequencing. *Nat Commun.* 2021;12(1):7198.
25. Lagarde J, et al. High-throughput annotation of full-length long noncoding RNAs with capture long-read sequencing. *Nat Genet.* 2017;49(12):1731–40.
26. Ståhl PL, et al. Visualization and analysis of gene expression in tissue sections by Spatial transcriptomics. *Science.* 2016;353(6294):78–82.
27. Rodrigues SG, et al. Slide-seq: A scalable technology for measuring genome-wide expression at high Spatial resolution. *Science.* 2019;363(6434):1463–7.
28. Asp M, Bergenstråhle J, Lundeberg J. Spatially resolved transcriptomes—next generation tools for tissue exploration. *BioEssays.* 2020;42(10):1900221.
29. Larsson L, Frisén J, Lundeberg J. Spatially resolved transcriptomics adds a new dimension to genomics. *Nat Methods.* 2021;18(1):15–8.
30. Mou T, et al. Reproducibility of methods to detect differentially expressed genes from single-cell RNA sequencing. *Front Genet.* 2020;10:1331.
31. Li Y, et al. Exaggerated false positives by popular differential expression methods when analyzing human population samples. *Genome Biol.* 2022;23(1):79.
32. Kim TK. T test as a parametric statistic. *Korean J Anesthesiology.* 2015;68(6):540–6.

Received: 18 August 2025 / Accepted: 25 November 2025

Published online: 06 December 2025

33. Law CW, et al. Voom: precision weights unlock linear model analysis tools for RNA-seq read counts. *Genome Biol.* 2014;15(2):R29.
34. Ritchie ME, et al. Limma powers differential expression analyses for RNA-sequencing and microarray studies. *Nucleic Acids Res.* 2015;43(7):e47–47.
35. Seyednasrollah F, Laiho A, Elo LL. Comparison of software packages for detecting differential expression in RNA-seq studies. *Brief Bioinform.* 2015;16(1):139–70.
36. Smyth G et al. Shi... W: Limma: linear models for microarray and RNA-Seq Data User's Guide. 2002, Citeseer.
37. Benjamini Y, Hochberg Y. Controlling the false discovery rate: a practical and powerful approach to multiple testing. *J Roy Stat Soc: Ser B (Methodol).* 1995;57(1):289–300.
38. Robinson MD, McCarthy DJ, Smyth GK. EdgeR: a bioconductor package for differential expression analysis of digital gene expression data. *Bioinformatics.* 2010;26(1):139–40.
39. Love M, Anders S, Huber W. Differential analysis of count data—the DESeq2 package. *Genome Biol.* 2014;15(550):10–1186.
40. Love MI, Huber W, Anders S. Moderated Estimation of fold change and dispersion for RNA-seq data with DESeq2. *Genome Biol.* 2014;15(12):550.
41. Colaprico A, et al. TCGAAbiolinks: an R/Bioconductor package for integrative analysis of TCGA data. *Nucleic Acids Res.* 2016;44(8):e71–71.
42. Mounir M, et al. New functionalities in the TCGAAbiolinks package for the study and integration of cancer data from GDC and GTEx. *PLoS Comput Biol.* 2019;15(3):e1006701.
43. Tarazona S, et al. Data quality aware analysis of differential expression in RNA-seq with NOISeq R/Bioc package. *Nucleic Acids Res.* 2015;43(21):e140–140.
44. Tarazona S, et al. Differential expression in RNA-seq: a matter of depth. *Genome Res.* 2011;21(12):2213–23.
45. Li J, Tibshirani R. Finding consistent patterns: a nonparametric approach for identifying differential expression in RNA-Seq data. *Stat Methods Med Res.* 2013;22(5):519–36.
46. Zhu A, et al. Nonparametric expression analysis using Inferential replicate counts. *Nucleic Acids Res.* 2019;47(18):e105–105.
47. Muzellec B, et al. PyDESeq2: a python package for bulk RNA-seq differential expression analysis. *Bioinformatics.* 2023;39(9):btad547.
48. Kharchenko PV, Silberstein L, Scadden DT. Bayesian approach to single-cell differential expression analysis. *Nat Methods.* 2014;11(7):740–2.
49. Finak G, et al. MAST: a flexible statistical framework for assessing transcriptional changes and characterizing heterogeneity in single-cell RNA sequencing data. *Genome Biol.* 2015;16(1):278.
50. Trapnell C, et al. Pseudo-temporal ordering of individual cells reveals dynamics and regulators of cell fate decisions. *Nat Biotechnol.* 2014;32(4):381.
51. Miao Z, et al. Dsingle for detecting three types of differential expression in single-cell RNA-seq data. *Bioinformatics.* 2018;34(18):3223–4.
52. Stuart T, et al. Comprehensive integration of single-cell data. *Cell.* 2019;177(7):1888–902. e21.
53. Hao Y, et al. Integrated analysis of multimodal single-cell data. *Cell.* 2021;184(13):3573–87. e29.
54. Qiu X, et al. Single-cell mRNA quantification and differential analysis with census. *Nat Methods.* 2017;14(3):309–15.
55. Van den Berge K, et al. Trajectory-based differential expression analysis for single-cell sequencing data. *Nat Commun.* 2020;11(1):1201.
56. Satija R, et al. Spatial reconstruction of single-cell gene expression data. *Nat Biotechnol.* 2015;33(5):495–502.
57. Butler A, et al. Integrating single-cell transcriptomic data across different conditions, technologies, and species. *Nat Biotechnol.* 2018;36(5):411–20.
58. Wolf FA, Angerer P, Theis FJ. SCANPY: large-scale single-cell gene expression data analysis. *Genome Biol.* 2018;19(1):15.
59. Gayoso A, et al. A python library for probabilistic analysis of single-cell omics data. *Nat Biotechnol.* 2022;40(2):163–6.
60. Zhang Y, et al. Scedar: A scalable python package for single-cell RNA-seq exploratory data analysis. *PLoS Comput Biol.* 2020;16(4):e1007794.
61. theislab/diffxpy. <https://github.com/theislab/diffxpy>.
62. Guo M, et al. SINCERA: a pipeline for single-cell RNA-Seq profiling analysis. *PLoS Comput Biol.* 2015;11(11):e1004575.
63. Nabavi S, et al. EMDomics: a robust and powerful method for the identification of genes differentially expressed between heterogeneous classes. *Bioinformatics.* 2016;32(4):533–41.
64. Soneson C, Robinson MD. Bias, robustness and scalability in single-cell differential expression analysis. *Nat Methods.* 2018;15(4):255–61.
65. Dal Molin A, Baruzzo G, Di Camillo B. Single-cell RNA-sequencing: assessment of differential expression analysis methods. *Front Genet.* 2017;8:62.
66. Aran D, et al. Reference-based analysis of lung single-cell sequencing reveals a transitional profibrotic macrophage. *Nat Immunol.* 2019;20(2):163–72.
67. Aibar S, et al. SCENIC: single-cell regulatory network inference and clustering. *Nat Methods.* 2017;14(11):1083–6.
68. Street K, et al. Slingshot: cell lineage and pseudotime inference for single-cell transcriptomics. *BMC Genomics.* 2018;19(1):477.
69. Cabello-Aguilar S, et al. SingleCellSignalR: inference of intercellular networks from single-cell transcriptomics. *Nucleic Acids Res.* 2020;48(10):e55–55.
70. Li C, et al. Multi-omic single-cell velocity models epigenome–transcriptome interactions and improves cell fate prediction. *Nat Biotechnol.* 2023;41(3):387–98.
71. Zhang K, et al. A fast, scalable and versatile tool for analysis of single-cell omics data. *Nat Methods.* 2024;21(2):217–27.
72. Huizing G-J, et al. Paired single-cell multi-omics data integration with mowgli. *Nat Commun.* 2023;14(1):7711.
73. Li B, et al. Cumulus provides cloud-based data analysis for large-scale single-cell and single-nucleus RNA-seq. *Nat Methods.* 2020;17(8):793–8.
74. Dorado. Documentation. <https://dorado-docs.readthedocs.io/en/latest/basecaller/mods/>.
75. Esfahani NG et al. Evaluation of nanopore direct RNA sequencing updates for modification detection. *bioRxiv*, 2025: p. 2025.05. 01.651717.
76. *Nanopore sequencing accuracy. Oxford Nanopore Technologies.* Available from: <https://nanoporetech.com/platform/accuracy>
77. Gleeson J, et al. Accurate expression quantification from nanopore direct RNA sequencing with nanocount. *Nucleic Acids Res.* 2022;50(4):e19–19.
78. Kuo RI, et al. Illuminating the dark side of the human transcriptome with long read transcript sequencing. *BMC Genomics.* 2020;21(1):751.
79. BrooksLabUCSC/flair. <https://github.com/BrooksLabUCSC/flair>.
80. Wyman D, et al. A technology-agnostic long-read analysis pipeline for transcriptome discovery and quantification. 2019. 10.1101/672931 Wyman D et al. A technology-agnostic long-read analysis pipeline for transcriptome discovery and quantification. *Biorxiv*. 2019: p. 672931. <https://doi.org/10.1101/672931>
81. Kovaka S, et al. Transcriptome assembly from long-read RNA-seq alignments with StringTie2. *Genome Biol.* 2019;20(1):278.
82. Pribelski AD, et al. Accurate isoform discovery with isoquant using long reads. *Nat Biotechnol.* 2023;41(7):915–8.
83. Differential Gene Expression Tutorial. https://epi2me.nanoporetech.com/notebooks/Differential_gene_expression.html.
84. epi2me-labs/wf-transcriptomes. <https://github.com/epi2me-labs/wf-transcriptomes>.
85. Marx V. Method of the year: spatially resolved transcriptomics. *Nat Methods.* 2021;18(1):9–14.
86. Moncada R, et al. Integrating microarray-based Spatial transcriptomics and single-cell RNA-seq reveals tissue architecture in pancreatic ductal adenocarcinomas. *Nat Biotechnol.* 2020;38(3):333–42.
87. A spatiotemporal organ-wide gene expression and cell atlas of the developing human heart. [https://www.cell.com/cell/fulltext/S0092-8674\(19\)31282-6](https://www.cell.com/cell/fulltext/S0092-8674(19)31282-6).
88. Svensson V, Teichmann SA, Stegle O. SpatialDE: identification of spatially variable genes. *Nat Methods.* 2018;15(5):343–6.
89. Sun S, Zhu J, Zhou X. Statistical analysis of Spatial expression patterns for Spatially resolved transcriptomic studies. *Nat Methods.* 2020;17(2):193–200.
90. Zhao E, et al. Spatial transcriptomics at Subspot resolution with bayesspace. *Nat Biotechnol.* 2021;39(11):1375–84.
91. Dries R, et al. Giotto: a toolbox for integrative analysis and visualization of Spatial expression data. *Genome Biol.* 2021;22(1):78.
92. Yuan Z, et al. Benchmarking Spatial clustering methods with Spatially resolved transcriptomics data. *Nat Methods.* 2024;21(4):712–22.
93. You Y, et al. Systematic comparison of sequencing-based Spatial transcriptomic methods. *Nat Methods.* 2024;21(9):1743–54.
94. Johnson WE, Li C, Rabinovic A. Adjusting batch effects in microarray expression data using empirical Bayes methods. *Biostatistics.* 2007;8(1):118–27.
95. Han Y, et al. Advanced applications of RNA sequencing and challenges. *Bioinform Biol Insights.* 2015;9:BBI.
96. Ozsolak F, Milos PM. RNA sequencing: advances, challenges and opportunities. *Nat Rev Genet.* 2011;12(2):87–98.
97. Tung P-Y, et al. Batch effects and the effective design of single-cell gene expression studies. *Sci Rep.* 2017;7(1):39921.
98. Vallejos CA, Richardson S, Marioni JC. Beyond comparisons of means: Understanding changes in gene expression at the single-cell level. *Genome Biol.* 2016;17(1):70.

99. Lun L, Bach ATK, Marioni JC. Pooling across cells to normalize single-cell RNA sequencing data with many zero counts. *Genome Biol.* 2016;17(1):75.
100. Tran HTN, et al. A benchmark of batch-effect correction methods for single-cell RNA sequencing data. *Genome Biol.* 2020;21(1):12.
101. Andrews TS, Hemberg M. Identifying cell populations with ScRNASeq. *Mol Aspects Med.* 2018;59:114–22.
102. Liu-Wei W, et al. Sequencing accuracy and systematic errors of nanopore direct RNA sequencing. *BMC Genomics.* 2024;25(1):528.
103. Wang Y, et al. Nanopore sequencing technology, bioinformatics and applications. *Nat Biotechnol.* 2021;39(11):1348–65.
104. Furlan M, et al. Computational methods for RNA modification detection from nanopore direct RNA sequencing data. *RNA Biol.* 2021;18(sup1):31–40.
105. Tzadikario T et al. Genomic in vitro transcription and nanopore direct RNA sequencing of a human B-Lymphocyte cell line. *bioRxiv*, 2025: p. 2025.04.25.650674.
106. Atta L, Fan J. Computational challenges and opportunities in spatially resolved transcriptomic data analysis. *Nat Commun.* 2021;12(1):5283.
107. Choe K, et al. Advances and challenges in Spatial transcriptomics for developmental biology. *Biomolecules.* 2023;13(1):156.
108. Du J, et al. Advances in Spatial transcriptomics and related data analysis strategies. *J Translational Med.* 2023;21(1):330.
109. Fang S, et al. Computational approaches and challenges in Spatial transcriptomics. *Genom Proteom Bioinform.* 2023;21(1):24–47.
110. Zhao Y, et al. Stereo-seq V2: Spatial mapping of total RNA on FFPE sections with high resolution. *Cell.* 2025; 188(23):6554–6571.e21Zhao Y et al. Stereo-seq V2: Spatial mapping of total RNA on FFPE sections with high resolution. *Cell.* 2025. <https://doi.org/10.1016/j.cell.2025.08.008>
111. Ren P, et al. Systematic benchmarking of high-throughput subcellular Spatial transcriptomics platforms across human tumors. *Nat Commun.* 2025;16(1):9232.
112. He S, et al. High-plex imaging of RNA and proteins at subcellular resolution in fixed tissue by Spatial molecular imaging. *Nat Biotechnol.* 2022;40(12):1794–806.

Publisher's note

Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.