

TECHNICAL NOTE

A single-cell RNA-sequencing training and analysis suite using the Galaxy framework

Mehmet Tekman ^{1,*†}, Bérénice Batut ^{1,†}, Alexander Ostrovsky ², Christophe Antoniewski ^{3,4}, Dave Clements ², Fidel Ramirez ⁵, Graham J. Etherington ⁶, Hans-Rudolf Hotz ^{7,8}, Jelle Scholtalbers ⁹, Jonathan R. Manning ¹⁰, Lea Bellenger ³, Maria A. Doyle ^{11,12}, Mohammad Heydarian ², Ni Huang ^{10,13}, Nicola Soranzo ⁶, Pablo Moreno ¹⁰, Stefan Mautner ¹, Irene Papatheodorou ¹⁰, Anton Nekrutenko ¹⁴, James Taylor ^{2,†}, Daniel Blankenberg ¹⁵, Rolf Backofen ¹ and Björn Grüning ^{1,*}

¹Department of Bioinformatics, University of Freiburg, Georges-Köhler-Allee 106, 79110 Freiburg, Germany;

²Department of Biology, Johns Hopkins University, Mudd Hall 144, 3400 N. Charles Street, Baltimore, MD 21218, USA; ³ARTbio, Sorbonne Université, CNRS FR 3631, Inserm US 037, Paris, France; ⁴Institut de Biologie Paris Seine, 9 Quai Saint-Bernard Université Pierre et Marie Curie, Campus Jussieu, Bâtiments A-B-C, 75005 Paris, France; ⁵Boehringer Ingelheim International GmbH, Binger Strasse 173, 55216 Ingelheim am Rhein, Biberach, Germany; ⁶Earlham Institute, Norwich Research Park, Norwich NR4 7UZ, UK; ⁷Friedrich Miescher Institute for Biomedical Research, Maulbeerstrasse 66, 4058 Basel, Switzerland; ⁸SIB Swiss Institute of Bioinformatics, Maulbeerstrasse 66, 4058 Basel, Switzerland; ⁹European Molecular Biology Laboratory, Genome Biology Unit, Meyerhofstraße 1, 69117 Heidelberg, Germany; ¹⁰European Molecular Biology Laboratory, European Bioinformatics Institute, EMBL-EBI, Wellcome Genome Campus, Hinxton, Cambridgeshire, CB10 1SD, UK; ¹¹Research Computing Facility, Peter MacCallum Cancer Centre, Melbourne, 305 Grattan Street, Victoria 3000, Australia; ¹²Sir Peter MacCallum Department of Oncology, The University of Melbourne, Victoria 3010, Australia; ¹³Wellcome Sanger Institute, Wellcome Genome Campus, Hinxton, Cambridgeshire, CB10 1SA, UK; ¹⁴Department of Biochemistry and Molecular Biology, The Pennsylvania State University, University Park, PA 16802, USA and ¹⁵Genomic Medicine Institute, Lerner Research Institute, Cleveland Clinic, 9500 Euclid Avenue, NB21 Cleveland, OH 44195, USA

*Correspondence address. Mehmet Tekman, Department of Bioinformatics, University of Freiburg, Georges-Köhler-Allee 106, 79110 Freiburg, Germany. E-mail: tekman@informatik.uni-freiburg.de  <http://orcid.org/0000-0002-4181-2676> and Björn Grüning, Department of Bioinformatics, University of Freiburg, Georges-Köhler-Allee 106, 79110 Freiburg, Germany. E-mail: gruening@informatik.uni-freiburg.de  <http://orcid.org/0000-0002-3079-6586>

[†]These authors contributed equally to this work.

[‡]Deceased.

Received: 22 June 2020; Revised: 30 August 2020

© The Author(s) 2020. Published by Oxford University Press GigaScience. This is an Open Access article distributed under the terms of the Creative Commons Attribution License (<http://creativecommons.org/licenses/by/4.0/>), which permits unrestricted reuse, distribution, and reproduction in any medium, provided the original work is properly cited.

Abstract

Background: The vast ecosystem of single-cell RNA-sequencing tools has until recently been plagued by an excess of diverging analysis strategies, inconsistent file formats, and compatibility issues between different software suites. The uptake of 10x Genomics datasets has begun to calm this diversity, and the bioinformatics community leans once more towards the large computing requirements and the statistically driven methods needed to process and understand these ever-growing datasets. **Results:** Here we outline several Galaxy workflows and learning resources for single-cell RNA-sequencing, with the aim of providing a comprehensive analysis environment paired with a thorough user learning experience that bridges the knowledge gap between the computational methods and the underlying cell biology. The Galaxy reproducible bioinformatics framework provides tools, workflows, and trainings that not only enable users to perform 1-click 10x preprocessing but also empower them to demultiplex raw sequencing from custom tagged and full-length sequencing protocols. The downstream analysis supports a range of high-quality interoperable suites separated into common stages of analysis: inspection, filtering, normalization, confounder removal, and clustering. The teaching resources cover concepts from computer science to cell biology. Access to all resources is provided at the singlecell.usegalaxy.eu portal. **Conclusions:** The reproducible and training-oriented Galaxy framework provides a sustainable high-performance computing environment for users to run flexible analyses on both 10x and alternative platforms. The tutorials from the Galaxy Training Network along with the frequent training workshops hosted by the Galaxy community provide a means for users to learn, publish, and teach single-cell RNA-sequencing analysis.

Keywords: scRNA; Galaxy; resources; high-performance computing; single-cell; 10x; training; Web

Background

Single-cell RNA sequencing and cellular heterogeneity

The continuing rise in single-cell technologies has led to previously unprecedented levels of analysis into cell heterogeneity within tissue samples, providing new insights into developmental and differentiation pathways for a wide range of disciplines. Gene expression studies are now performed at a cellular level of resolution, which, compared to bulk RNA-sequencing (RNA-seq) methods, allows researchers to model their tissue samples as distributions of different expressions instead of as an average.

Pathways from single-cell data

The various expression profiles uncovered within tissue samples infer discrete cell types that are related to one another across an “expression landscape.” The relationships between the more distinct profiles are inferred via distance-metrics or manifold learning techniques. Ultimately, the aim is to model the continuous biological process of cell differentiation from multipotent stem cells to distinct mature cell types, and infer lineage and differentiation pathways between transient cell types [1].

Elucidating cell identity

Trajectory analysis that integrates the up- or downregulation of significant genes along lineage branches can then be performed to uncover the factors and extracellular triggers that can coerce a pluripotent cell to become biased towards one cell fate outcome compared to another. This undertaking has created a new frontier of exploration in cell biology, where researchers have assembled reference maps for different cell lines for the purpose of fully recording these cell dynamics and their characteristics to create a global “atlas” of cells [2, 3].

Pitfalls and technical challenges Sequencing sensitivity and normalization. With each new protocol comes a host of new technical problems to overcome. The first wave of software utilities to deal with the analysis of single-cell datasets were statistical

packages, aimed at tackling the issue of “dropout events” during sequencing, which would manifest as a high prevalence of zero-entries in >80% of the feature count matrix. These zeroes were problematic because they could not be trivially ignored as their presence stated either that the cell did not produce any molecules for that transcript or that the sequencer simply did not detect them. Normalization techniques originally developed for bulk RNA-seq had to be adapted to accommodate for this uncertainty, and new ones were created that harnessed hurdle models, data imputation via manifold learning techniques, or by pooling subsets of cells together and building general linear models [4].

Improvements in sequencing. With the downstream analysis packages attempting to solve the dropouts via stochastic methods, the upstream sequencing technologies also aspired to solve the capture efficiency via new well, droplet, and flow cytometry-based protocols, all of which lend importance to the process of barcoding sequencing reads.

In each protocol, cells are tagged with cell barcodes such that any reads derived from them can be unambiguously assigned to the cell of origin. Unique molecular identifiers (UMIs) are also used to mitigate the effects of amplification bias of transcripts within the same cell. The detection, extraction, and (de-)multiplexing of cell barcodes and UMIs is therefore one of the first hurdles that researchers encounter when receiving raw FASTQ data from a sequencing facility.

The burgeoning software ecosystem

Since its conception, several different packages and many pipelines have been developed to assist researchers in the analysis of single-cell RNA sequencing (scRNA-seq) [5, 6]. Most of these packages were written for the R programming language because many of the novel normalization methods developed to handle the dropout events depended on statistical packages that were primarily R-based [7]. Stand-alone analysis suites emerged as the different authors of these packages rapidly expanded their methods to encapsulate all facets of the single-cell analysis, often creating compatibility issues with previous package versions. The Bioconductor repository provided some much-needed stability in this regard by hosting stable releases, but re-

searchers were still prone to building directly from repository sources in order to reap the benefits of new features in the upstream versions [8, 9].

Nonexchangeable data formats. Another issue was the proliferation of the many different and quickly evolving R-based file formats for processing and storing the data, such as SingleCellExperiment as used by the Scater suite, SCSeq from RaceID, and SeuratObject from Seurat [10, 11]. Many new packages would cater only to one format or suite, leading to the common problem that data processed in one suite could not be reliably processed by methods in another. This incompatibility between packages fuelled a choice of one analysis suite over another, or conversely required researchers to dig deeper into the internal semantics of R S4 objects to manually slot data components together [12]. These problems only accelerated the rapid development of these suites, leading to further version instability. As a result of this analysis diversity, there are many tutorials on how to perform scRNA-seq analysis, each oriented around one of these pipelines [13].

Error propagation and analysis uncertainty. A crucial quality control step upstream, such as filtering or the removal of unwanted variability, can propagate forward into the downstream sections to yield wildly different results on the same data. This uncertainty, and the statistically driven methods to overcome it, leaves a wide knowledge gap for researchers simply trying to understand the underlying dynamics of cell identity.

Rise of 10x Genomics

10x launch. In 2015, 10x Genomics released their GemCode product, which was a droplet-sequencing-based protocol capable of sequencing tens of thousands of cells with an average cell quality higher than other facilities [14]. This unprecedented level of throughput steadily gained traction amongst researchers and start-ups seeking to perform single-cell analysis, and thus 10x datasets began to prevail in the field.

10x analysis software. 10x Genomics provided software that was able to perform much of the pre-processing, and provided feature-count matrices in a transparent HDF5-based format that provided a means of efficient matrix storage and exchange, and conclusively removed the restriction for downstream analysis modules to be written in R.

ScanPy, a popular alternative. The ScanPy suite [15], written in Python using its own HDF5-based AnnData format, became a valid alternative for analysing 10x datasets. The Seurat developers had similar aspirations and soon adopted the LOOM format, another HDF5 variant. However, the popularity of ScanPy increased as it began to integrate the methods of other standalone packages into its codebase, making it the natural choice for users who wanted to achieve more without compromising on compatibility between different suites [9].

Solutions in the cloud

Accessible science. As the size of datasets scaled, so did the computing resources required to analyse them, in terms of both processing power and storage. Galaxy is an open source biocomputing infrastructure that exemplifies the 3 main tenets of science: reproducibility, peer review, and open access—all freely accessible within the web browser [16]. It hosts a wide range of highly cited bioinformatics tools with many different versions and enables users to freely create their own workflows via a seamless drag-and-drop interface.

Reproducible workflows. Galaxy can make use of Conda or Containers to set up tool environments to ensure that the bioinformatics tools will always be able to run, even when the library

dependencies for a tool have changed, by building tools under locked version dependencies and bundling them together in a self-contained environment [17]. These environments provide a concise solution for the package version instability that plagues scRNA-seq analysis notebooks, in terms of both reproducibility and analysis flexibility. A user could keep the quality control results obtained from an older version of ScanPy whilst running a newer ScanPy version at the clustering stage to reap the benefits of the later improvements in that algorithm. By allowing the user to select from multiple versions of the same tool, and by further permitting different versions of the tools within a workflow, Galaxy enables an unprecedented level of free-flow analysis by letting researchers pick and choose the best aspects of a tool without worrying about the underlying software libraries [18]. The burdens of software incompatibility and choice of programming language that plagued the scRNA-seq analysis ecosystem before are now completely removed from the user.

User-driven custom workflows. Analyses are not limited to one pipeline either because the datasets that are passed between tools can easily be interpreted by a different tool that is capable of reading that dataset. In the case of scRNA-seq, Galaxy can convert between CSV, MTX, LOOM, and AnnData formats. This interexchange of modules from different tools further extends the flexibility of the analysis by again letting the user decide which component of a tool would be best suited for a specific part of an analysis.

Training resources. Galaxy also provides a wide range of learning resources, with the aim of guiding users step by step through an analysis, often reproducing the results of published works. The teaching and training materials are part of the Galaxy Training Network (GTN), which is a worldwide collaborative effort to produce high-quality teaching materials to educate users in how to analyse their data, and in turn to train others on the same materials via easily deployable workshops backed by monthly stable releases of the GTN materials [19]. Training materials are provided on a wide variety of different topics, and workshops are hosted regularly, as advertised on the Galaxy Events web portal. The GTN has grown rapidly since its conception and gains new volunteers every year, who each contribute and coordinate training and teaching events, maintain topics and subtopics, translate tutorials into multiple languages, and provide peer review on new material [20].

Methods

Stable workflows in Galaxy

The analysis of scRNA-seq within Galaxy was a 2-pronged effort concentrated on bringing high-quality single-cell tools into Galaxy and providing the necessary workflows and training to accompany them. As mentioned in the previous section, this effort needed to overcome incompatible file format issues and unstable packages due to rapid development and needed to establish a standardized basis for the analysis.

Tutorials

The tutorials are split into 2 main parts as outlined in Fig. 1: first, the pre-processing stage, which constructs a count matrix from the initial sequencing data; second, the cluster-based downstream analysis on the count matrix. These stages are very different from one another in terms of their information content: the pre-processing stage requires the researcher to be more familiar with wetlab sequencing protocols than a typical bioin-

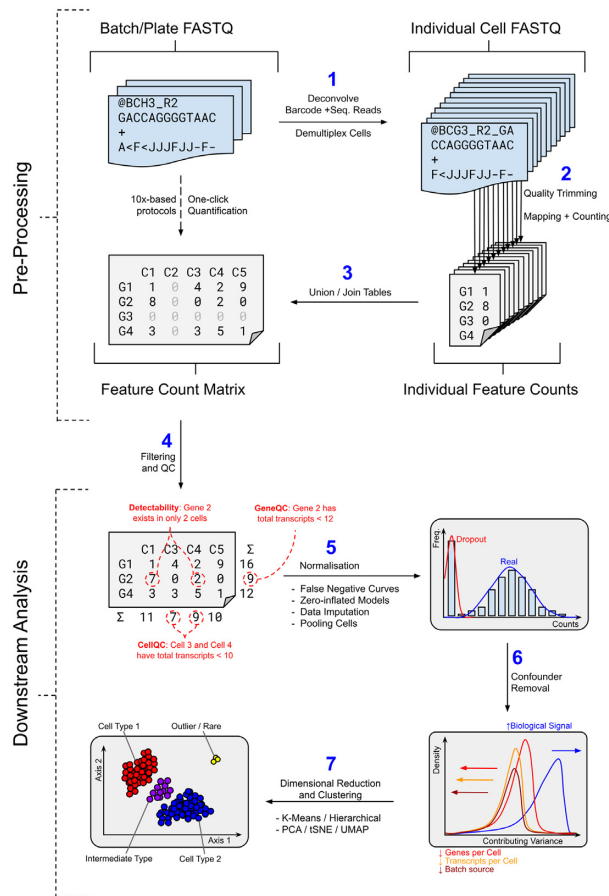


Figure 1: The main stages of single-cell analysis, separated broadly into the upper and lower stages of pre-processing and downstream analysis, respectively. Top: The 2 main routes to generating a count matrix from sequencing data: via 1-click quantification solutions or through manual demultiplexing. Bottom: The 4 main stages required to perform cluster-based analysis from the count matrix, through filtering, normalization, confounder removal, and embedding.

formatician would normally know, and the downstream analysis stage requires the researcher to be familiar with statistics concepts that a wetlab scientist might not be too familiar with. The tutorials are designed to broadly appeal to both the biologist and the statistician, as well as complete beginners to the entire topic.

Pre-processing workflows

The pre-processing scRNA-seq materials tackle the 2 most common use-cases that researchers will encounter when they first enter the field: processing scRNA-seq data from 10x Genomics, and processing data generated from alternative protocols. For instance, microwell-based protocols have been known to yield more features and display lower levels of dropouts compared to 10x, and so we accommodate for them by providing a more customizable path through the pre-processing stage [21].

Barcode extraction. Before the era of 10x Genomics, scRNA-seq data had to be demultiplexed, mapped, and quantified. The demultiplexing stage entails an intimate knowledge of cell barcodes and UMIs, which are protocol dependent, and expects the bioinformatician to know exactly where and how the data were generated. One common pitfall at this very first stage is estimating how many cells to expect from the FASTQ input data, and this requires 3 crucial pieces of information: which reads contain

the barcodes (or precisely, which subset of both the forward and reverse reads contains the barcodes); of these barcodes, which specific ones were actually used for the analysis; and how to re-solve barcode mismatches/errors.

Barcode estimation. Naive strategies involve using a known barcode template and querying against the FASTQ data to profile the number of reads that align to a specific barcode, often employing “knee” methods to estimate this amount [22]. However, this approach is not robust; certain cells are more likely to be overrepresented compared to others, and some cell barcodes may contain more unmappable reads compared to others, meaning that the metric of higher library read depth is not necessarily correlated with a better-defined cell. Ultimately, the bioinformatician must inquire directly with the sequencing laboratory as to which cell barcodes were used, because these are often not specific to the protocol but to the technician who designed them, with the idea that they should not align to a specific reference genome or transcriptome.

One-click pre-processing

Quantification with Cell Ranger. 10x Genomics simplified the single-cell RNA package ecosystem by using a language-independent file format and streamlining much of the barcode particularities with their Cell Ranger pipeline, allowing researchers to focus more on the internal biological variability of their datasets [23].

Quantification with STARsolo. The pre-processing workflow (titled “10X StarSolo Workflow”) in Galaxy uses the RNA STARsolo utility as a drop-in replacement for Cell Ranger, not only because it is a feature update of the already existing RNA STAR tool in Galaxy but because it boasts a 10-fold speedup in comparison to Cell Ranger and does not require Illumina lane-read information to perform the processing [24, 25].

Other approaches. The pre-processing workflows for these “1-click” solutions consume the same datasets and yield approximately the same count matrices by following similar modes of barcode discovery and quantification. Within Galaxy, there is also Alevin (paired with Salmon) and scPipe, which can both also perform the necessary demultiplexing, (alignment-free) mapping, and quantification stages in a single step [26, 27, 28].

Flexible pre-processing

CELSeq2 barcoding. The custom pre-processing workflow (titled “CELSeq2: Single Batch mm10”) is modelled after the CEL-seq2 protocol using the barcoding strategies of the Freiburg Max-Planck Institute laboratory as its main template, but the workflow is actually flexible to accommodate any droplet- or well-based protocol such as SMART-seq2 and Drop-seq [29].

Manual demultiplexing and quantification. The training pictographically guides users through the concepts of extracting cell barcodes from the protocol, explains the significance of UMIs in the process of read deduplication with illustrative examples, and instructs the user in the process of performing further quality controls on their data during the post-mapping process via RNA STAR and other tools that are native to Galaxy.

Training the user. At each stage, the user’s knowledge is queried via question prompts and expandable answer box dialogs, as well as helpful hints for future processing in comment boxes, all written in the transparent Markdown specification developed for contributing to the GTN.

Downstream workflows

Common stages of analysis. The downstream modules are defined by the 5 main stages of downstream scRNA-seq analysis:

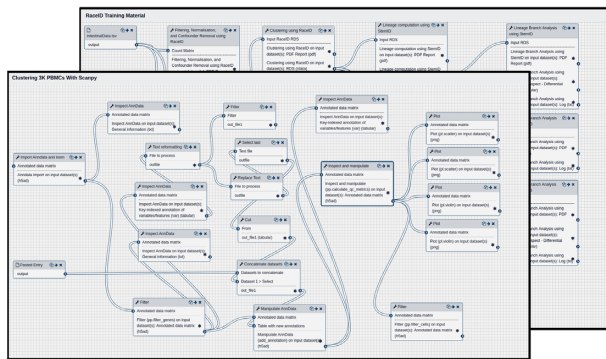


Figure 2: Downstream analysis workflows as shown in the Galaxy Workflow Editor for (top) RaceID and (bottom) ScanPy, each displaying modules symbolizing the 5 main stages of analysis.

sis: filtering, normalization, confounder removal, clustering, and trajectory inference. There are 3 workflows to aid in this process (2 of which are shown in Fig. 2), each sporting a different well-established scRNA-seq pipeline tool.

Quality control with Scater. The Scater pipeline follows a visualize-filter-visualize paradigm, which provides an intuitive means to perform quality control on a count matrix by use of repeated incremental changes on a dataset through the use of principal component analysis and library size-based metrics [30]. Once this pre-analysis stage is complete, the full downstream analysis (comprising the 5 stages mentioned above) can be performed by workflows based on the following suites: RaceID and ScanPy.

Downstream analysis with the RaceID suite. RaceID was developed initially to analyse rare cell transcriptomes whilst being robust against noise and thus is ideal for working with smaller datasets in the range of 300–1,000 cells. Owing to its complex cell lineage and fate prediction models, it can also be used on larger datasets with some scaling costs.

Downstream analysis with the ScanPy suite. ScanPy was developed as the Python alternative to the innumerable packages for scRNA-seq that were based on R, which was the dominant language for such analyses, and it was one of the first packages with native 10x Genomics support. Since then it has grown substantially and has been re-implementing much of the newer R-based methods released in BioConductor as “recipe” modules, thereby providing a single source to perform many different types of the same analysis.

The workflows derived from both these suites emulate the 5 main stages of analysis mentioned previously, where filtering, normalization, and confounder removal are typically separated into distinct stages, though sometimes merged into 1 step depending on the tool.

Filtering

Cell and gene removal. During the filtering stage, the initial count matrix removes low-quality or unwanted cells using typical parameters such as minimum gene detection sensitivity and minimum library size, and low-quality genes are also removed under similar metrics, where the minimum number of cells for a gene to be included is decided. The Scater pre-analysis workflow can also be used here to provide a principal component analysis-based method of feature selection so that only the highly variable genes are left in the analysis.

Disadvantages of filtering. There is always the danger of over-filtering a dataset, whereby setting overzealous lower-bound

thresholds on gene variability can have the undesired effect of removing essential housekeeping genes. These relatively uniformly expressed genes are often required for setting a baseline to establish a threshold to distinguish the more desired differentially expressed genes. It is therefore important that the user first perform a naive analysis and only later refine their filtering thresholds to boost the biological signal.

Normalization

Library size normalization. The normalization step aims to remove any technical factors that are not relevant to the analysis, such as the library size, where cells sharing the same identity are likely to differ from one another more by the number of transcripts they exhibit than due to more relevant biological factors.

Intrinsic cell factors. The first and foremost is cell capture efficiency, where different cells produce more or fewer transcripts based on the amplification and coverage conditions in which they are sequenced. The second is the presence of dropout events, which manifest as a prevalence of “zeroes” in the final count matrix. Whether a “zero” is imputable to nondetection of an existing molecule or to the absence of the molecule in the cell is uncertain. This uncertainty alone has led to a wide selection of different normalization techniques that try to model this expression either via hurdle models, imputing the data via manifold learning techniques, or working around the issue by pooling subsets of cells together [31].

In this regard, both the RaceID and ScanPy workflows offer many different normalization techniques, and users are encouraged to take advantage of the branching workflow model of Galaxy to explore all possible options.

Confounder removal

Regression of cell cycle effects. Other sources of variability stem from unwanted biological contributions known as confounder effects, such as cell cycle effects and transcription. Depending on the stage of the cell cycle at which a cell was sequenced, 2 cells of the same type might cluster differently because one might have more transcripts as a result of it being in the M-phase of the cell cycle. Library sizes notwithstanding, it is the variability in specific cell cycle genes that can be the main driving factor in the overall variability. Thankfully, these effects are easy to regress out, and we replicate an entire stand-alone ScanPy workflow dedicated to detecting and visualizing the effects based on the original notebook [32].

Transcriptional bursting. The transcription effects are harder to model because these are semi-stochastic and remain not well understood. In bulk RNA-seq the expression of genes undergoing transcription is averaged to give “high” or “low” signals, producing a global effect that gives the false impression that transcription is a continuous process. The reality is more complex, where cells undergo transcription in “bursts” of activity followed by periods of no activity, at irregular intervals [33]. At the bulk level these discrete processes are smoothed to give a continuous effect, but at the cell level it could mean that even 2 directly adjacent cells of the same type normalized to the same number of transcripts can still have different levels of expression for a gene due to this process. This is not something that can be countered for, but it does educate the users in which factors they can or cannot control in an analysis and how much variability they can expect to see.

Clustering and projection

Dimension reduction and clustering. Once a user has obtained a count matrix in which they are confident, they are then guided

through the process of dimension reduction (with choice of different distance metrics), choosing a suitable low-dimensional embedding, and performing clustering through commonly used techniques such as *k*-means, hierarchical, and neighbourhood community detection. These complex techniques are illustrated in layman's terms through the use of helpful images and community examples. For example, the GTN ScanPy tutorial explains the Louvain clustering approach [34] via a stand-alone slide deck to assist in the workflow [35].

Commonly used embeddings. The clustering and the cluster inspection stages are notably separated into distinct utilities here, with the understanding that the same initial clustering can appear dissimilar under different projections, e.g., *t*-distributed stochastic network embedding (*t*-SNE) against Uniform Manifold Approximation and Projection (UMAP) [36, 37]. Ultimately the user is encouraged to play with the plotting parameters to yield the best-looking clusters.

Static plots or interactive environments. Cluster inspection tools are available that allow users to easily generate static plots tailored to pipeline-specific information as originally defined by the software package authors. However, the AnnData and LOOM specifications store these map projection data separately, enabling the use of a plethora of possible plotting tools, including HTML5-based interactive visualizations, such as cellxgene [38], that permit on-demand querying and rendering of individual cell features without the need to generate static images. A collection of these Galaxy interactive tools can be accessed at the website live.usegalaxy.eu. Although these tools are excellent at dynamically displaying map projections, especially 3D ones, further computation must be performed to complete a full pseudotime analysis.

Pseudotime trajectory analysis

Inferring developmental pathways. The cell pseudotime series analysis is often referred to as the trajectory inference stage because cells are ordered along a trajectory to reflect the continuous changes of gene expression along a developmental pathway under the assumption that the cells are transitioning from one pluripotent type to another less potent type.

Pseudotime techniques. For the trajectory inference stage, there is the Partition-based Graph Abstraction (PAGA) technique championed by ScanPy [39], and there are also the FateID and StemID packages for the RaceID workflow [40]. The former provides a level of graph abstraction to the datasets in order to infer a community graph structure, which it can use to learn the shape of the data and infer pathways between neighbourhoods. The latter is more intuitive, in that it constructs a minimal spanning tree of related clusters that infer lineage, and cell fate decisions that can be explored by querying branches in the tree, as a function of the genes that are up- or downregulated along the currently explored pathway. The statistical strength and significance of each pathway guides the user along more valid trajectories that would more accurately reflect the biological variation occurring within transitioning cells.

Sharing reference maps

The insights and novel cell types discovered in these analyses can also be integrated into the Human Cell Atlas portal [41], which is an initiative that aims to classify unique or rare cell types, as well as their transitive properties, in order to build a comprehensive map of cells that can be used to investigate the various differentiation pathways of multipotent stem cells in the human body.

Galaxy Training Network

Tutorial Hierarchy

Tutorials in the GTN are grouped by topic, e.g., Variant Calling, Transcriptomics, Assembly. These tutorials can also declare prerequisites, so that users can review required concepts from previous tutorials, e.g., quality control checks from bulk RNA-seq still being used in scRNA-seq. Not only does this allow users to derive a clear route through the range of training materials, but it also empowers them to choose their own learning path through the network of topics. In particular for scRNA-seq, users can start their training from pre-processing tutorials and continue till downstream analysis.

Tutorial structure

Tutorials usually consist of a hands-on workflow that guides the user through an analysis with Galaxy utilizing a step-by-step approach, often accompanied by a slide deck that either serves to explain stand-alone concepts more concisely or is used during workshops and trainings as a way to introduce the user to the topic. In an effort to maintain reproducibility in science, all tutorials require example workflows, and all materials needed to run the workflows and tutorials are hosted for free with open access at Zenodo with a permanent DOI tag.

User-driven contribution

The user contributions are the heart of the GTN community, and options are given to appeal to different levels of contribution. At the casual level, each tutorial has at the bottom an anonymous feedback form that rates the quality of the tutorial and asks for hints on what could be improved, which the tutorial authors can then act on. At the more eager level, users can contribute directly to the material hosted at the GitHub repository using the approachable GTN Markdown format, which further empowers contributors not only to adapt existing material but also to write tutorials in their own specialist topics. The GitHub code reviews, paired with the plain-text GTN Markdown format, facilitate easy peer review of tutorial topics by using standard diff utilities.

Galaxy subdomains and environments

Subdomains encapsulate relevant tools

The Galaxy tools and the GTN are further tied together by Galaxy subdomains, which better serve the various topics within their own self-contained environments. These complement the training materials by providing only the necessary Galaxy tools in order to not trouble the user with unrelated tools that might not be so relevant to the material; e.g., Variant Analysis tools are not included in an scRNA-seq environment. This also has the benefit that smaller more specialized Galaxy instances can be packaged and deployed, avoiding the overhead of presenting the entire Galaxy tool repertoire.

Single Cell and Human Cell Atlas

In this light, the singlecell.usegalaxy.eu subdomain hosts the entirety of the single-cell materials, tools, workflows, and single-cell-related events. The full list of tools in the subdomain, as well as their application to the aforementioned stages of scRNA-seq analysis, is detailed in Supplemental Table 1. Human Cell Atlas community members, led by the European Bioinformatics Institute and the Wellcome Sanger Institute, have their own subdomain at humancellatlas.usegalaxy.eu [42], providing access to widely applicable tools including ScanPy, Seurat, and Monocle3

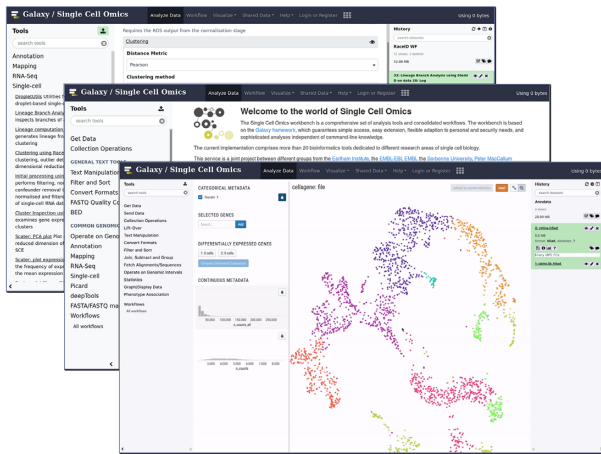


Figure 3: Galaxy Training Network hosting a comprehensive suite of tools, trainings, and workflows to perform scRNA-seq analysis.

[43], but also specialist tools such as those for cell type prediction (including scmap [44], scPred [45], and Garnett [46]).

Analysis in Galaxy workflows

The tools outlined in the Downstream workflows subsection expose the full set of parameters of their underlying program suites in order to serve the same level of customization that the users would expect when running a notebook-based analysis. This suits the needs of most researchers, but some are more used to processing the data directly in a language-driven notebook environment.

Galaxy Interactive Environments

For the more computer programming-oriented users, Galaxy hosts interactive environments at live.usegalaxy.eu, which allows users to spin up their own Jupyter [47] or RStudio [48] notebooks whilst harnessing the same cloud compute infrastructure. Here, users can import their Galaxy datasets, process them in their own desired manner, and export them back into their histories in a similar way to how datasets are treated in workflows.

List of Galaxy Interactive Environments

In addition to interactive notebooks, the Galaxy Interactive Environments also boasts a selection of other interactive tools such as the aforementioned cellxgene featured in Fig. 3, as well as SPARQL, a query language interface; BAM/VCF IOBIO, a file format analysis viewer [49]; EtherCalc, a web spreadsheet [50]; PHINCH, a metagenome visualizer [51]; Wallace, a species modelling platform [52]; WILSON, an omics visualizer [53]; IDE for materials science; Panoply, a netCDF viewer [54]; HiGlass, a Hi-C data visualizer [55]; and even an XFCE Virtual Desktop environment [56].

Discussion

Growth of scRNA training materials

The amount of single-cell materials on the GTN is increasing substantially every year, with at first only 1 pre-processing tutorial in 2018, 1 downstream tutorial at the start of 2019, and at the present time of writing 3 pre-processing tutorials and 3 downstream analysis workflows, further accompanied by slide decks and interactive visualizations. Single-cell Galaxy workshops based on these materials have been given at the Single-

Cell RNAseq Training Course 2018 at the Earlham Institute, the 2019 Galaxy Community Conference, within the Freiburg Meln-Bio Consortium, and at the Association of Biomolecular Resource Facilities. The trainings also lend themselves seamlessly to online webinars, which have proved useful during the COVID-19 lockdown period.

Reproducible cloud-based analysis

The advent of scRNA-seq analysis within the Galaxy framework re-echoes the efforts to standardize the analysis of scRNA-seq, with the promise of presenting reproducible research. The burden of computation on the ever-growing size of the datasets is shifted to the cloud computing resources, and as scRNA-seq technology scales, more researchers are likely to migrate towards cloud-based solutions to reap the benefits of superior computing abilities and storage capabilities. Ultimately, the Galaxy framework abstracts the user from the many non-trivial technicalities of the analysis and exposes them to a legible interface of tools from which they can pick and choose.

Longevity and accessibility

The community regularly comes together during scheduled code festivals (CoFests) or hackathons to review, contribute, and actively maintain the training materials. The number of community contributions has steadily increased over the past 4 years [16], and this growth trend ensures that the Galaxy resources will stay current and adapt to changes in scRNA-seq technology and analysis methods if necessary. The GTN also makes use of language translation tools to provide international interpretations of the training materials to reach a wider more internationally diverse audience.

Future of scRNA-seq in Galaxy

The capacity for growth of scRNA-seq in Galaxy is limitless, with the continuing acquisition of new single-cell tools being incorporated into Galaxy workflows, and the expanding GTN community bringing more expert-level contributions to the training material. The vestiges of incompatible libraries and inexchangeable file formats are unburdened from the user as the epoch of web-based tools and strong biocomputing frameworks becomes more dominant. From the first data upload to the final finishing touches of a customized workflow, the single-cell Galaxy portal upholds the ideals of open science by supporting the user all the way from the initial training to the final publication, where they can export and share their results with a single click.

Availability of Source Code and Requirements

- Project name: Single-Cell RNA-seq Analysis in Galaxy
- Project home page: singlecell.usegalaxy.eu
- Operating system(s): Web-based, platform independent
- License: GNU GPL v3

Availability of Supporting Data and Materials

All datasets used in the GTN are independently hosted at Zenodo and are easily findable under the tag “Galaxy Training Network,” as well as being directly hosted within the Galaxy Data Libraries on the UseGalaxy.eu server.

The tool wrappers that serve as the functional components of the many different single-cell analysis tools are hosted at the

GitHub Tools-IUC repository, as well as at the Galaxy Toolshed under the category of “Transcriptomics.”

Additional Files

Supplemental Table 1. Comprehensive table of all single-cell RNA-analysis related software packages within Galaxy and their capabilities.

Abbreviations

DOI: Digital Object Identifier; GTN: Galaxy Training Network; HDF5: Hierarchical Data Format 5; PAGA: Partition-based Graph Abstraction; RNA-seq: RNA sequencing; scRNA-seq: single-cell RNA sequencing; t-SNE: t-distributed stochastic network embeddings; UMAP: Uniform Manifold Approximation and Projection; UMI: Unique Molecular Identifier.

Competing Interests

The authors declare that they have no competing interests.

Funding

This project was funded by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation) 22977937/GRK2344 and BA2168/3-3, Germany's Excellence Strategy (CIBSS - EXC-2189 - Project ID 390939984), the BBSRC Core Strategic Programme Grants BBS/E/T/000PR9814, BBS/E/T/000PR9817, BBS/E/T/000PR9818, and BBS/E/T/000PR9819, Core Capability Grant BBS/E/T/000PR9816 at the Earlham Institute, and the National Institutes of Health grant U41HG006620.

The European Galaxy project is in part funded by Collaborative Research Centre 992 Medical Epigenetics (DFG grant SFB 992/1 2012) and German Federal Ministry of Education and Research (BMBF grants 031 A538A/A538C RBC, 031L0101B/031L0101C de.NBI-epi). The article processing charge was funded by the Baden-Württemberg Ministry of Science, Research and Art and the University of Freiburg in the funding programme Open Access Publishing.

Authors' Contributions

B.G. conceived the project and created the singlecell subdomain. M.T. and B.B. wrapped RaceID and ScanPy, respectively, and created all initial workflows and trainings. A.O., B.G., C.A., D.C., D.B., G.J.E., H.R.H., J.R.M., L.B., M.A.D., M.H., N.H., F.R., J.S., N.S., P.M., and S.M. developed tools or made functional contributions to the tools and training materials. D.B., I.P., A.N., J.T., and R.B. supported the development of the project. M.T. wrote the original draft manuscript. All authors have read, made suggestions, and ultimately approved the final manuscript.

Acknowledgements

We thank the bioinformatics group at the University of Freiburg for the development and hosting of the European Galaxy server, Monika Degen-Hellmuth at the Backofen Lab for her assistance in the organization of the project, the Institut Français de Bioinformatique (IFB) for its support of the ARTbio team, Charles Girardot at EMBL Heidelberg for his useful feedback, and we also acknowledge the worldwide contributions from users and de-

velopers towards the Galaxy Project and all upstream authors and contributors of the software ecosystem that we use and rely on.

References

- Wagner A, Regev A, Yosef N. Revealing the vectors of cellular identity with single-cell genomics. *Nat Biotechnol* 2016;**34**(11):1145.
- Rozenblatt-Rosen O, Stubbington MJ, Regev A, et al. The Human Cell Atlas: from vision to reality. *Nature* 2017;**550**(7677):451–3.
- Briggs JA, Weinreb C, Wagner DE, et al. The dynamics of gene expression in vertebrate embryogenesis at single-cell resolution. *Science* 2018;**360**(6392):eaar5780.
- Camara PG. Methods and challenges in the analysis of single-cell RNA-sequencing data. *Curr Opin Syst Biol* 2018;**7**:47–53.
- McCarthy DJ, Campbell KR, Lun AT, et al. Scater: pre-processing, quality control, normalization and visualization of single-cell RNA-seq data in R. *Bioinformatics* 2017;**33**(8):1179–86.
- Satija R, Farrell JA, Gennert D, et al. Spatial reconstruction of single-cell gene expression data. *Nat Biotechnol* 2015;**33**(5):495.
- Amezquita RA, Lun AT, Becht E, et al. Orchestrating single-cell analysis with Bioconductor. *Nat Methods* 2020;**17**(2):137–45.
- Satija R, Farrell JA, Gennert D, et al. List of Seurat Releases. <https://github.com/satijalab/seurat/releases>. Accessed 10 January 2020.
- Wolf F, Angerer P, Theis F. ScanPy Release Notes. <http://scanpy.readthedocs.io/en/stable/release-notes.html>. Accessed 10 January 2020.
- Lun A, Risso D, Korthauer K. 2018. SingleCellExperiment: S4 classes for single cell data. R package version 1.10.1, doi:10.18129/B9.bioc.SingleCellExperiment.
- Grün D, Lyubimova A, Kester L, et al. Single-cell messenger RNA sequencing reveals rare intestinal cell types. *Nature* 2015;**525**(7568):251.
- Ruckdeschel P, Kohl M, Stabla T, et al. S4 classes for distributions. *R News* 2006;**6**(2):2–6.
- Luecken MD, Theis FJ. Current best practices in single-cell RNA-seq analysis: a tutorial. *Mol Syst Biol* 2019;**15**(6):e8746.
- Vickovic S, Ståhl PL, Salmén F, et al. Massive and parallel expression profiling using microarrayed single-cell sequencing. *Nat Commun* 2016;**7**:13182.
- Wolf FA, Angerer P, Theis FJ. SCANPY: large-scale single-cell gene expression data analysis. *Genome Biol* 2018;**19**(1):15.
- Afgan E, Baker D, Batut B, et al. The Galaxy platform for accessible, reproducible and collaborative biomedical analyses: 2018 update. *Nucleic Acids Res* 2018;**46**(W1):W537–44.
- Grüning B, Chilton J, Köster J, et al. Practical computational reproducibility in the life sciences. *Cell Syst* 2018;**6**(6):631–5.
- Grüning B, Dale R, Sjödin A, et al. Bioconda: sustainable and comprehensive software distribution for the life sciences. *Nat Methods* 2018;**15**(7):475.
- Batut B, Hiltmann S, Bagnacani A, et al. List of Galaxy Training Network Releases. <https://github.com/galaxyproject/training-material/releases>. Accessed 10 January 2020.
- Batut B, Hiltmann S, Bagnacani A, et al. Community-driven data analysis training for biology. *Cell Syst* 2018;**6**(6):752–8.

21. Wang X, Yao H, Zhang Q, et al. Direct comparative analysis of 10X Genomics Chromium and Smart-seq2. *bioRxiv* 2019:615013.
22. Smith T, Heger A, Sudbery I. UMI-tools: modeling sequencing errors in Unique Molecular Identifiers to improve quantification accuracy. *Genome Res* 2017;27(3):491–9.
23. Zheng GX, Terry JM, Belgrader P, et al. Massively parallel digital transcriptional profiling of single cells. *Nat Commun* 2017;8(1):14049.
24. Dobin A, Davis CA, Schlesinger F, et al. STAR: ultrafast universal RNA-seq aligner. *Bioinformatics* 2013;29(1):15–21.
25. Dobin A. STARsolo Release Page. <https://github.com/alexdobin/STAR/blob/master/docs/STARsolo.md>. Accessed 10 January 2020.
26. Srivastava A, Malik L, Smith T, et al. Alevin efficiently estimates accurate gene abundances from dscRNA-seq data. *Genome Biol* 2019;20(1):65.
27. Patro R, Duggal G, Love MI, et al. Salmon provides fast and bias-aware quantification of transcript expression. *Nat Methods* 2017;14(4):417–9.
28. Tian L, Su S, Dong X, et al. scPipe: a flexible R/Bioconductor preprocessing pipeline for single-cell RNA-sequencing data. *PLoS Comput Biol* 2018;14(8):e1006361.
29. Hashimshony T, Senderovich N, Avital G, et al. CEL-Seq2: sensitive highly-multiplexed single-cell RNA-Seq. *Genome Biol* 2016;17(1):77.
30. Etherington GJ, Soranzo N, Mohammed S, et al. A Galaxy-based training resource for single-cell RNA-sequencing quality control and analyses. *GigaScience* 2019;8(12):giz144.
31. Lun AT, Bach K, Marioni JC. Pooling across cells to normalize single-cell RNA sequencing data with many zero counts. *Genome Biol* 2016;17(1):75.
32. Wolf F, Angerer P, Rybakov S. ScanPy Preprocessing and Clustering 3k PBMCs Tutorial. <https://scanpy-tutorials.readthedocs.io/en/latest/pbmc3k.html>. Accessed 10 January 2020.
33. Raj A, van Oudenaarden A. Nature, nurture, or chance: stochastic gene expression and its consequences. *Cell* 2008;135(2):216–26.
34. Blondel VD, Guillaume JL, Lambiotte R, et al. Fast unfolding of communities in large networks. *J Stat Mech* 2008;2008(10):P10008.
35. Tekman M. Accompanying Slide Deck for ScanPy PBMC Workflow. <https://training.galaxyproject.org/training-material/topics/transcriptomics/tutorials/scrna-scanpy-pbmc3k/slides.html>. Accessed 10 January 2020.
36. Maaten Lvd, Hinton G. Visualizing data using t-SNE. *J Mach Learn Res* 2008;9:2579–05.
37. McInnes L, Healy J, Melville J. Umap: Uniform manifold approximation and projection for dimension reduction. *arXiv* 2018:1802.03426.
38. McGill C, Weaver C, Martin B et al., chanzuckerberg/cellxgene: Release 0.11.2. Zenodo 2019, doi:10.5281/zenodo.3368662.
39. Wolf FA, Hamey FK, Plass M, et al. PAGA: graph abstraction reconciles clustering with trajectory inference through a topology preserving map of single cells. *Genome Biol* 2019;20(1):59.
40. Herman JS, Sagar , Grün D. FateID infers cell fate bias in multipotent progenitors from single-cell RNA-seq data. *Nat Methods* 2018;15(5):379–86.
41. Regev A, Teichmann SA, Lander ES, et al. Science Forum: The Human Cell Atlas. *Elife* 2017;6:e27041.
42. Moreno P, Huang N, Manning JR, et al. User-friendly, scalable tools and workflows for single-cell analysis. *bioRxiv* 2020, doi:10.1101/2020.04.08.032698.
43. Cao J, Spielmann M, Qiu X, et al. The single-cell transcriptional landscape of mammalian organogenesis. *Nature* 2019;566(7745):496–502.
44. Kiselev VY, Yiu A, Hemberg M. scmap: projection of single-cell RNA-seq data across data sets. *Nat Methods* 2018;15(5):359–62.
45. Alquicira-Hernandez J, Sathe A, Ji HP, et al. scPred: accurate supervised method for cell-type classification from single-cell RNA-seq data. *Genome Biol* 2019;20(1):264.
46. Pliner HA, Shendure J, Trapnell C. Supervised classification enables rapid annotation of cell atlases. *Nat Methods* 2019;16(10):983–6.
47. Kluyver T, Ragan-Kelley B, Pérez F, et al. Jupyter Notebooks—a publishing format for reproducible computational workflows. In: Loizides F, Schmidt B , eds. Positioning and Power in Academic Publishing: Players, Agents and Agendas. IOS; 2016:87–90.
48. Allaire J. RStudio: integrated development environment for R. Boston, MA; 2012:770.
49. Miller C, Qiao Y, DiSera T, et al. Bam. Iobio: a Web-based, real-time, sequence alignment file inspector. *Nat Methods* 2014;11:1189.
50. Tang A. EtherCalc Github Repository. <https://github.com/audreyt/ethercalc>. Accessed 10 January 2020.
51. Bik HM, Pitch Interactive Inc. Phinch: an interactive, exploratory data visualization framework for –Omic datasets. *bioRxiv* 2014, doi:10.1101/009944.
52. Kass JM, Vilela B, Aiello-Lammens ME, et al. Wallace: a flexible platform for reproducible modeling of species niches and distributions built for community expansion. *Methods Ecol Evol* 2018;9(4):1151–6.
53. Schultheis H, Kuenne C, Preussner J, et al. WillsON: Web-based Interactive Omics Visualization. *Bioinformatics* 2019;35(6):1055–7.
54. Schmunk RB. Panoply netcdf, hdf and grib data viewer. National Aeronautics and Space Administration-Goddard Institute for Space Studies, 2018.
55. Kerpedjiev P, Abdennur N, Lekschas F, et al. HiGlass: web-based visual exploration and analysis of genome interaction maps. *Genome Biol* 2018;19(1):125.
56. Fourdan O. Xfce: a lightweight desktop environment. Annual Linux Showcase and Conference, Atlanta, usenix.org;2000:1–6.