

scPNMF: sparse gene encoding of single cells to facilitate gene selection for targeted gene profiling

Dongyuan Song^{1,†}, Kexin Li^{2,†}, Zachary Hemminger^{3,4}, Roy Wollman^{3,4,5} and Jingyi Jessica Li^{2,6,7,8,*}

¹Bioinformatics Interdepartmental Ph.D. Program, University of California, Los Angeles, CA 90095-7246, USA, ²Department of Statistics, University of California, Los Angeles, CA 90095-1554, USA, ³Institute for Quantitative and Computational Biosciences, University of California, Los Angeles, CA 90095, USA, ⁴Department of Integrative Biology and Physiology, University of California, Los Angeles, CA 90095-7239, USA, ⁵Department of Chemistry and Biochemistry, University of California, Los Angeles, CA 90095-1569, USA, ⁶Department of Human Genetics, University of California, Los Angeles, CA 90095-7088, USA, ⁷Department of Computational Medicine, University of California, Los Angeles, CA 90095-1766, USA and ⁸Department of Biostatistics, University of California Los Angeles, CA 90095-1772, USA

*To whom correspondence should be addressed.

Abstract

Motivation: Single-cell RNA sequencing (scRNA-seq) captures whole transcriptome information of individual cells. While scRNA-seq measures thousands of genes, researchers are often interested in only dozens to hundreds of genes for a closer study. Then, a question is how to select those informative genes from scRNA-seq data. Moreover, single-cell targeted gene profiling technologies are gaining popularity for their low costs, high sensitivity and extra (e.g. spatial) information; however, they typically can only measure up to a few hundred genes. Then another challenging question is how to select genes for targeted gene profiling based on existing scRNA-seq data.

Results: Here, we develop the single-cell Projective Non-negative Matrix Factorization (scPNMF) method to select informative genes from scRNA-seq data in an unsupervised way. Compared with existing gene selection methods, scPNMF has two advantages. First, its selected informative genes can better distinguish cell types. Second, it enables the alignment of new targeted gene profiling data with reference data in a low-dimensional space to facilitate the prediction of cell types in the new data. Technically, scPNMF modifies the PNMF algorithm for gene selection by changing the initialization and adding a basis selection step, which selects informative bases to distinguish cell types. We demonstrate that scPNMF outperforms the state-of-the-art gene selection methods on diverse scRNA-seq datasets. Moreover, we show that scPNMF can guide the design of targeted gene profiling experiments and the cell-type annotation on targeted gene profiling data.

Availability and implementation: The R package is open-access and available at <https://github.com/JSB-UCLA/scPNMF>. The data used in this work are available at Zenodo: <https://doi.org/10.5281/zenodo.4797997>.

Supplementary information: [Supplementary data](#) are available at *Bioinformatics* online.

Contact: jli@stat.ucla.edu

1 Introduction

The recent development of single-cell RNA sequencing (scRNA-seq) technologies provides unprecedented opportunities to decipher transcriptome heterogeneity among individual cells (Birnbau, 2018; Potter, 2018; Zhu *et al.*, 2020). A typical scRNA-seq dataset contains thousands to tens of thousands of genes; however, a subset of genes, which we call *informative genes*, are usually sufficient for representing the underlying biological variations of cells in the dataset for two reasons. First, variations of many genes are not related to the biological variations of interest. For instance, fluctuations in the expression levels of housekeeping genes are irrelevant to cell types (Eisenberg and Levanon, 2013; Thellin *et al.*, 1999). Second, many genes have strongly correlated expression levels, suggesting that one gene may represent a group of genes without much loss of information (Subramanian *et al.*, 2017). Therefore, for scRNA-seq data

analysis, informative gene selection has three advantages: (i) enhancing biological signals by removing unwanted technical variations, (ii) improving the interpretability of analysis results by focusing on informative genes and (iii) reducing the number of genes to save computational resources.

Besides scRNA-seq data analysis, informative gene selection is also crucial for designing single-cell targeted gene profiling experiments, which we define to include all technologies that measure only a specific set of genes' expression levels in individual cells. Unlike scRNA-seq, targeted gene profiling requires a limited number (often no more than hundreds) of genes to be specified before sequencing. Examples of targeted gene profiling include spatial technologies [e.g. smFISH (Raj *et al.*, 2008) and MERFISH (Moffitt *et al.*, 2016)] and non-spatial technologies [e.g. BART-Seq (Uzbas

et al., 2019), HyPR-seq (Marshall *et al.*, 2020) and 10x-Genomics Targeted Gene Expression]. Compared with scRNA-seq, targeted gene profiling technologies have advantages such as capturing spatial information (by smFISH and MERFISH), having a lower cost per cell (by BART-Seq), and exhibiting a higher sensitivity for detecting lowly expressed genes (by HyPR-seq). However, it remains an open and challenging question to optimize the gene selection for targeted gene profiling under a gene number limitation.

Given the importance of informative gene selection, researchers have developed many gene selection methods for scRNA-seq data. Most existing methods select genes based on the relationship between per-gene expression means and per-gene expression variances (with the mean and variance of each gene calculated across cells). Popular example methods include variance stabilization transformation (vst) (Hafemeister and Satija, 2019) and mean-variance plot (mvp) in the R package Seurat (Stuart *et al.*, 2019), as well as modelGeneVar in the R package scan (Lun *et al.*, 2016). These methods select highly variable genes that have large expression variances in relation to their expression means. Other methods use various metrics of gene importance instead of the per-gene expression variance. For example, M3Drop selects the genes that have zero expression levels in many cells (Andrews and Hemberg, 2019); GiniClust selects the genes with large Gini indices of expression levels (Jiang *et al.*, 2016); SCMarker selects the genes that have expression levels bi/multimodally distributed and are coexpressed or mutually exclusively expressed with some other genes (Wang *et al.*, 2019). A common limitation of these existing methods is that they are all designed to select a relatively large number of genes; thus, their performance for selecting a small number of genes remains unclear. For instance, in Seurat, the default gene number is 2000; SCMarker selects 700–900 genes in its exemplar applications (Wang *et al.*, 2019). All these gene numbers are much greater than 200, the maximum gene number allowed by multiple targeted gene profiling technologies. Therefore, existing gene selection methods may not be suitable for selecting genes for targeted gene profiling. Another drawback of these methods is that their selected genes lack functional interpretability; that is, their selected genes are not categorized as functional gene groups.

In addition to these gene selection methods, linear dimensional-reduction methods, such as principal component analysis (PCA) and non-negative matrix factorization (NMF), can also be used for gene selection. Specifically, genes can be selected based on their contributions to the projected low dimensions found by PCA or NMF (Baron *et al.*, 2016; Macosko *et al.*, 2015; Ye and Li, 2016; Zhu *et al.*, 2017). Although many variants of PCA and NMF algorithms have been developed for scRNA-seq data analysis, they are not designed for gene selection (Boileau *et al.*, 2020; Duren *et al.*, 2018; Durif *et al.*, 2019; Gao and Welch, 2020; Welch *et al.*, 2019; Yang and Michailidis, 2016; Zhang *et al.*, 2020).

Here, we propose an unsupervised method scPNMF to simultaneously select informative genes and project scRNA-seq data onto an interpretable low-dimensional space. Leveraging the Projective Non-negative Matrix Factorization (PNMF) algorithm (Yuan *et al.*, 2009), scPNMF combines the advantages of PCA and NMF by outputting a non-negative sparse weight matrix that can project cells in a high-dimensional scRNA-seq dataset onto a low-dimensional space. Unlike the weight matrix (a.k.a., loading matrix) found by PCA, the non-negative sparse weight matrix output by scPNMF corresponds to bases that each corresponds to a group of coexpressed genes. Compared with the original PNMf, a unique feature of scPNMF is basis selection: scPNMF uses correlation screening and multimodality testing to remove the bases that cannot reveal potential cell clusters in the input scRNA-seq dataset. There are two functionalities of scPNMF: (i) given a prespecified gene number and a scRNA-seq dataset, scPNMF selects informative genes based on its weight matrix; (ii) given a targeted gene profiling dataset containing the informative genes, scPNMF projects this dataset onto the same low-dimensional space of a reference scRNA-seq dataset containing cell-type labels, thus enabling cell-type annotation on the targeted gene profiling dataset. A comprehensive benchmark shows that scPNMF outperforms existing gene selection methods in two

aspects. First, the informative genes selected by scPNMF lead to the most accurate cell clustering. Second, the informative genes and weight matrix of scPNMF lead to the best cell-type prediction accuracy for targeted gene profiling data. Therefore, scPNMF is a powerful gene selection method that can guide the experimental design and data analysis of single-cell targeted gene profiling.

2 Materials and methods

The core of scPNMF is to learn a low-dimensional embedding of cells so that the bases of the low-dimensional space correspond to sparse and mutually exclusive gene groups, and that genes in each group are coexpressed and thus functionally related. Figure 1 illustrates the workflow of scPNMF. The input of scPNMF is a log-transformed gene-by-cell count matrix measured by scRNA-seq. There are two main steps in scPNMF: (i) it learns a low-dimensional sparse weight matrix by PNMf; (ii) it selects bases in the weight matrix based on functional annotations (optional), correlation screening and multimodality testing to remove uninformative bases that cannot distinguish cell types. The output of scPNMF includes (i) the selected weight matrix, a sparse and mutually exclusive encoding of genes as new, low dimensions and (ii) the score matrix containing embeddings of input cells in the low dimensions. The selected weight matrix has two main applications: extracting informative genes for downstream analyses, such as cell clustering and new marker gene identification, and projecting new targeted gene profiling data for data integration and cell-type annotation.

2.1 scPNMF Step I: PNMf

In this section, we review the PNMf algorithm (Yuan *et al.*, 2009; Yang and Oja, 2010) as the foundation of scPNMF. We first compare the formulation of PNMf with that of PCA and NMF, and we show that PNMf has the advantages of both PCA and NMF so that it can be a useful tool for scRNA-seq data analysis. Next, we introduce our PNMf implementation.

Given a log-transformed count matrix $\mathbf{X} \in \mathbb{R}_{\geq 0}^{p \times n}$, whose p rows correspond to genes and whose n columns represent cells, and a positive integer $K \leq p$, PNMf aims to find a K -dimensional space, whose dimensions correspond to non-negative, sparse and mutually exclusive linear combinations of the p genes, so that projecting the n cells onto the K -dimensional space does not cause much information loss (i.e. projecting the K -dimensional embeddings of the n cells back to the original p -dimensional space can largely restore the original n cells). PNMf tackles this task by solving the optimization problem:

$$\min_{\mathbf{W} \in \mathbb{R}_{\geq 0}^{p \times K}} \|\mathbf{X} - \mathbf{W}\mathbf{W}^T\mathbf{X}\|, \quad (1)$$

where $\|\cdot\|$ denotes the Frobenius matrix norm. The solution $\mathbf{W}$ is referred to as a *weight matrix*. Each column of $\mathbf{W}$ is a *basis*, whose p entries are the weights of the p genes. PNMf requires all weights to be non-negative, leading to a sparse $\mathbf{W}$ with most weights as zeros.

PCA is similar to PNMf but does not require all weights to be non-negative. We can write the optimization problem of PCA as

$$\min_{\mathbf{W} \in \mathbb{R}^{p \times K}, \mathbf{W}^T\mathbf{W}=\mathbf{I}} \|\mathbf{X} - \mathbf{W}\mathbf{W}^T\mathbf{X}\|, \quad (2)$$

whose solution $\mathbf{W}$ is also a weight matrix but not sparse, and $\mathbf{W}$ is often referred to as the loading matrix.

A common property of PNMf and PCA is that the transpose of their weight matrix, $\mathbf{W}^T \in \mathbb{R}^{K \times p}$, can be used to project a new cell with p gene measurements, $\mathbf{x} \in \mathbb{R}^p$, onto the K -dimensional space as $\mathbf{W}^T\mathbf{x}$.

In contrast to PNMf and PCA, NMF finds two non-negative matrices $\mathbf{W}$ and $\mathbf{H}$ so that their product approximates the original matrix $\mathbf{X}$. NMF solves the optimization problem:

$$\min_{\mathbf{W} \in \mathbb{R}_{\geq 0}^{p \times K}, \mathbf{H} \in \mathbb{R}_{\geq 0}^{K \times n}} \|\mathbf{X} - \mathbf{W}\mathbf{H}\|, \quad (3)$$

whose solution $\mathbf{W}$ still has K columns representing bases, and $\mathbf{H}$ has

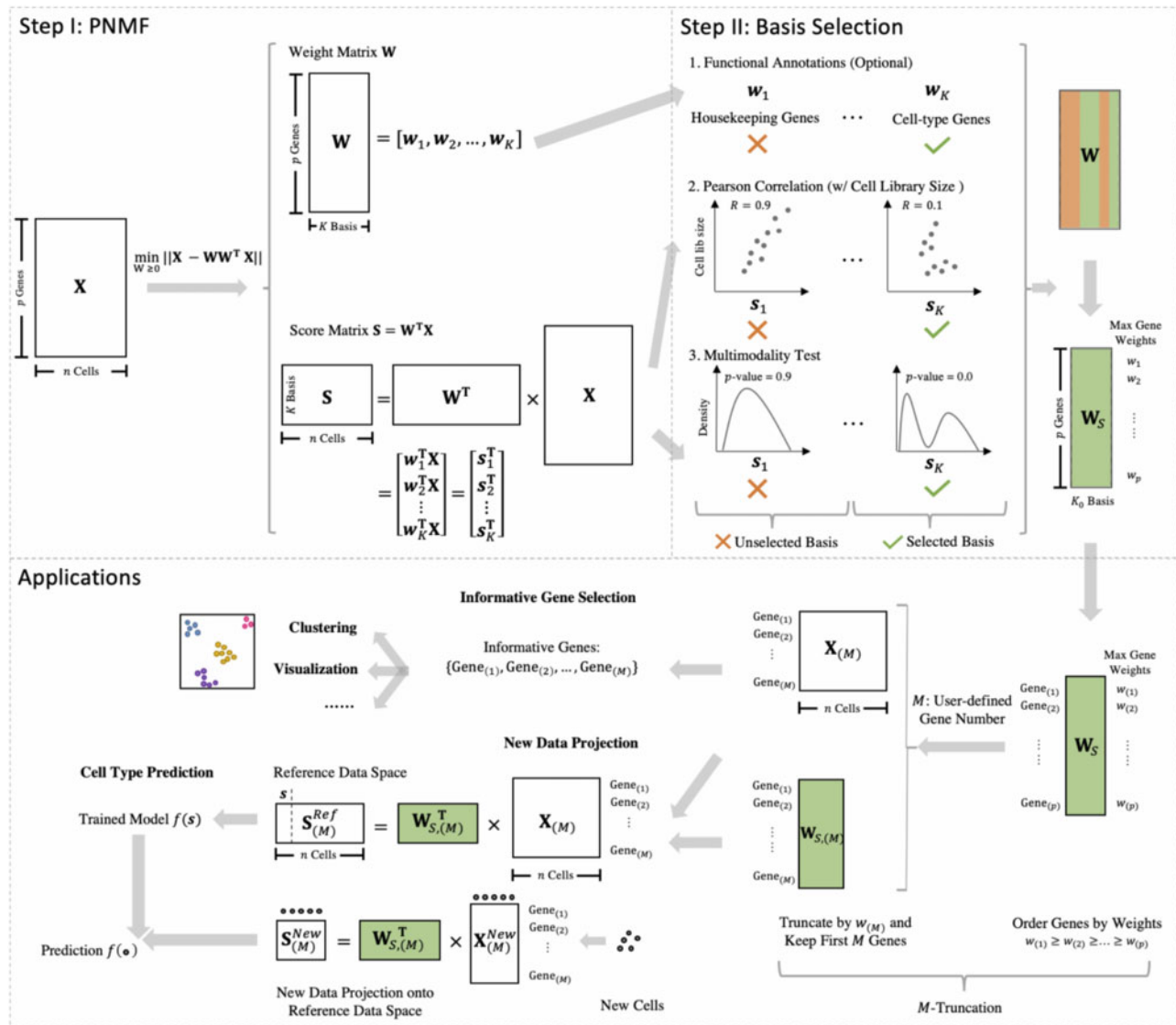


Fig. 1 An overview of scPNMF. Taking a log-transformed gene-by-cell count matrix as the input, scPNMF first learns a low-dimensional sparse weight matrix W and a low-dimensional cell embedding matrix S . Second, it removes the bases irrelevant to cell-type variations by examining bases' functional annotations (optional), Pearson correlations with cell library sizes, and multimodality. Given a user-defined gene number M , scPNMF performs M -truncation to facilitate two main applications: (1) selecting the desired number of informative genes; (2) projecting new targeted gene profiling data onto the low-dimensional space defined by reference scRNA-seq data. The details are in the Methods section.

n columns as K -dimensional embeddings of the n cells. Due to the non-negative constraint on W and H , W is a sparse matrix (Lee and Seung, 1999). However, the transpose W^T cannot be used as a projection matrix from the original p -dimensional space to a K -dimensional space. The reason is that, if W^T is a projection matrix, then by the definition of H we have $W^T X = H$, which would convert the objective function (3) of NMF to the objective function (1) of PNMF. In other words, PNMF is a constrained version of NMF by requiring W^T to be a projection matrix. Hence, PNMF inherits the property of NMF by having non-negative, sparse bases that are mostly mutually exclusive (i.e., different bases correspond to different gene groups). Moreover, based on the similarities of the objective functions of PNMF (1) and PCA (2), we can see that PNMF also resembles PCA by finding a weight matrix whose transpose can serve as a projection matrix and whose bases are largely orthogonal to each other. Table 1 summarizes the properties of PNMF, PCA and NMF.

In the context of scRNA-seq data analysis, the above advantages of PNMF lead to an interpretable and useful weight matrix W . First,

the high sparsity of W makes each basis (column) depend on only a small set of genes, which has been defined as a *meta-gene* for NMF (Brunet et al., 2004). Second, the mutual exclusiveness of W makes different bases correspond to different gene sets, easing the interpretation of bases as meta-genes or functional units. Third, the projection matrix W^T allows the alignment of new data to reference data, thus facilitating cell-type annotation on the new data. Algorithm 1 summarizes the key steps of PNMF implementation in scPNMF. Our implementation mainly follows the two papers that proposed the PNMF algorithm (Yang and Oja, 2010; Yuan et al., 2009), and we change the initialization of W to the weight matrix found by PCA, W_{PCA} , with the absolute value taken on every entry. Our initialization is motivated by the desired orthogonality of bases (i.e. columns of W).

With the weight matrix $W \in \mathbb{R}_{\geq 0}^{p \times K}$ learned by PNMF, we obtain the score matrix $S = W^T X \in \mathbb{R}_{\geq 0}^{K \times n}$, whose K rows correspond to the bases and whose n columns represent the cells. Specifically, the j th column of S is the K -dimensional embedding of the j th cell; the

Table 1 Comparison of the properties of PNMf, PCA and NMF

	Optimization problem	Non-negativity	Sparsity	Mutual exclusiveness	New data projection
PNMF	$\min_{\mathbf{W}} \ \mathbf{X} - \mathbf{W}\mathbf{W}^T\mathbf{X}\ \text{ s.t. } \mathbf{W} \geq 0$	Yes	Very high	Very high	Yes
PCA	$\min_{\mathbf{W}} \ \mathbf{X} - \mathbf{W}\mathbf{W}^T\mathbf{X}\ \text{ s.t. } \mathbf{W}^T\mathbf{W} = \mathbf{I}$	No	Low	Low	Yes
NMF	$\min_{\mathbf{W}, \mathbf{H}} \ \mathbf{X} - \mathbf{W}\mathbf{H}\ \text{ s.t. } \mathbf{W}, \mathbf{H} \geq 0$	Yes	High	High	No

Algorithm 1 Pseudocode of PNMf implementation in scPNMF

Initialize: $\mathbf{W} = \text{abs}(\mathbf{W}_{\text{PCA}}) \in \mathbb{R}_{\geq 0}^{p \times K}$
1: while not converge **do**
2: for $i = 1, \dots, p; k = 1, \dots, K$ **do**
3: $\mathbf{W}_{ik} \leftarrow \mathbf{W}_{ik} \frac{2(\mathbf{X}\mathbf{X}^T\mathbf{W})_{ik}}{(\mathbf{W}\mathbf{W}^T\mathbf{X}\mathbf{X}^T\mathbf{W})_{ik} + (\mathbf{X}\mathbf{X}^T\mathbf{W}\mathbf{W}^T\mathbf{X})_{ik}}$
4: end for
5: $\mathbf{W} \leftarrow \frac{1}{\|\mathbf{W}\|_2} \mathbf{W}$
6: end while
Output: $\mathbf{W} \in \mathbb{R}_{\geq 0}^{p \times K}, \mathbf{S} = \mathbf{W}^T\mathbf{X} \in \mathbb{R}_{\geq 0}^{K \times n}$

k th row of $\mathbf{S}$, denoted by $\mathbf{s}_k^T$, contains the *scores* (i.e. coordinates) of all n cells in the k th basis:

$$\mathbf{s}_k^T = \mathbf{w}_k^T \mathbf{X}, \quad (4)$$

where $\mathbf{w}_k$ is the k th column of $\mathbf{W}$, $k = 1, \dots, K$.

The low rank K needs to be prespecified in PNMf, same as in PCA and NMF. A larger K preserves more information in $\mathbf{X}$ but also removes less noise (technical variation of cells that is not of biological interest), impedes the interpretation of $\mathbf{W}$ (more bases are more difficult to interpret) and increases the computational burden. To choose K in a data-driven way, we propose an orthogonality measure, which shows that $K=20$ is a reasonable choice for multiple scRNA-seq datasets (Supplementary Section S1.1).

2.2 scPNMF Step II: basis selection

The second key step of scPNMF is to select informative bases among the K bases found by PNMf (i.e. columns of $\mathbf{W}$ and rows of $\mathbf{S}$) to remove unwanted variations of cells (e.g. variations irrelevant to cell types). The columns of $\mathbf{W}$ enjoy high sparsity and mutual exclusiveness; that is, each column contains positive weights corresponding to a unique small set of genes, so it is expected to reflect a certain biological function. However, some biological functions may not be relevant to the cell heterogeneity of interest, for example, cell-type composition. Motivated by this, we propose three strategies for selecting informative bases (columns of $\mathbf{W}$ and rows of $\mathbf{S}$): functional annotations (optional), correlations with cell library sizes and tests of multimodality.

2.2.1 Strategy 1: examine bases by functional annotations (optional)

The first, optional strategy is to annotate the biological function(s) of each basis in the weight matrix. For example, scPNMF may apply gene ontology (GO) analysis to the top 10% genes with the highest weights in each basis (column of $\mathbf{W}$) and record the enriched GO terms as the basis' functional annotation. Then, users with prior knowledge can interpret the functional annotation on each basis and decide whether or not to remove the basis. For example, if the goal is to delineate cell types in scRNA-seq data, a basis corresponding to cell-cycle genes should be removed because they would obscure the distinction of cell types.

However, it is worth noting that filtering bases by biological annotations is optional in scPNMF. Conservative users can keep all K bases output by PNMf and directly use data-driven basis selection (Section 2.2.2). For our results in this article, scPNMF removes the

bases corresponding to well-known housekeeping genes (Supplementary Section S2).

2.2.2 Data-driven strategies

2.2.2.1 Strategy 2: examine bases by correlations with cell library sizes. Note that the input of scPNMF is a log-transformed unnormalized count matrix for users' convenience. Hence, scPNMF does not adjust for cell library sizes in the computation of $\mathbf{W}$ and $\mathbf{S}$ in Step I. (For a detailed discussion on why scPNMF uses unnormalized data as input, see Supplementary Section S6.) Given that the variance of cell library sizes contributes to unwanted variations of cells (Hafemeister and Satija, 2019), it is necessary to remove the bases whose corresponding rows in $\mathbf{S}$ are strongly correlated with cell library sizes.

We use the total log-transformed counts to approximate the library size of each cell, and we calculate the Pearson correlation between each $\mathbf{s}_k$ and the library sizes of n cells. The strategy is to retain the bases whose Pearson correlations are under a pre-defined threshold, which we set to 0.7 based on empirical observations (Supplementary Section S1.2).

2.2.2.2 Strategy 3: examine bases by multimodality tests. Another data-driven strategy is to retain the bases whose corresponding scores are multimodally distributed. If a basis' score vector (row in $\mathbf{S}$) contains n scores with a multimodality pattern, then it is likely to distinguish cell types and should be retained. To implement this strategy, we use the ACR test (Ameijeiras-Alonso et al., 2019) to check the multimodality of each basis' score vector. The null hypothesis is that the score vector contains n scores sampled from a unimodal distribution, and the alternative hypothesis is that the distribution has more than one mode. After performing multiple multimodality tests, one per basis, we use the Benjamini-Hochberg procedure to set a P -value threshold by controlling the false discovery rate under 1%. The bases whose P -values are under this threshold will be retained.

In summary, scPNMF step II allows users to use Strategy 1 to filter out uninformative bases based on functional annotations if available; then it implements data-driven Strategies 2 and 3 to further remove bases that have strong correlations with cell library sizes and exhibit unimodality patterns. The retained bases will have their corresponding columns in $\mathbf{W}$ selected and stacked into the *selected weight matrix* $\mathbf{W}_S \in \mathbb{R}_{\geq 0}^{p \times K_0}$, where K_0 is the number of selected bases.

2.3 Applications of scPNMF output: informative gene selection and new data projection

The selected weight matrix $\mathbf{W}_S$ output by scPNMF has two main applications: selection of a desired number of informative genes and projection of new targeted gene profiling data onto the low-dimensional space defined by $\mathbf{W}_S$. Given a gene number M (e.g. 200), scPNMF uses M -truncation, a step to select M rows in $\mathbf{W}_S$, resulting in M informative genes and a *truncated, selected weight matrix* $\mathbf{W}_{S,(M)} \in \mathbb{R}_{\geq 0}^{M \times K_0}$ for new data projection.

2.3.1 M -Truncation and informative gene selection

We denote the desired number of informative genes by $M \in \mathbb{N}$, with $M \leq$ of nonzero rows in $\mathbf{W}_S$. M -truncation has three steps.

1. For each gene i , calculate its largest weight w_i across bases in $\mathbf{W}_S$:

$$w_i = \max_{k=1,\dots,K_0} (\mathbf{W}_S)_{ik}, \quad i = 1, 2, \dots, p. \quad (5)$$

2. Order genes by their maximum weights $w_{(1)} \geq w_{(2)} \geq \dots \geq w_{(p)}$ and set the truncation threshold as $w_{(M)}$. Identify the first M genes as *informative genes*.
3. Construct the truncated, selected weight matrix $\mathbf{W}_{S,(M)}$:
 - 3.1. Truncate the selected weight matrix $\mathbf{W}_S$ by setting all $(\mathbf{W}_S)_{ik} < w_{(M)}$ to be 0;
 - 3.2. Keep the M rows with nonzero entries; stack them by row into $\mathbf{W}_{S,(M)}$ based on the order of the informative genes.

In short, scPNMF selects informative genes based on their maximum weights in the selected bases. The rationale is that a gene's maximum weight reflects the gene's contribution to the establishment of the K_0 -dimensional space, which preserves the n cells' biological variations of interest. Hence, genes with larger maximum weights are more informative in the sense of encoding cells' biological variations. An important application of informative gene selection is to guide the design of targeted gene profiling experiments.

2.3.2 New data projection

Given the selected M informative genes, once new cells are measured by targeted gene profiling on these genes, $\mathbf{W}_{S,(M)}$ can be used to project the new cells onto the K_0 -dimensional space where the cells in the input scRNA-seq data are embedded in. If the input data has cell-type annotations, we refer to the input data as *reference data*, then we can predict the new cells' types from the types of the cells in the reference data. In detail, new data projection has the following steps:

1. Apply scPNMF with M -truncation to input, reference data $\mathbf{X} \in \mathbb{R}_{\geq 0}^{p \times n}$ with n cells to obtain the truncated, selected weight matrix $\mathbf{W}_{S,(M)}$. Construct $\mathbf{X}_{(M)} \in \mathbb{R}_{\geq 0}^{M \times n}$ as a submatrix of $\mathbf{X}$, with rows corresponding to the rows of $\mathbf{W}_{S,(M)}$, that is, the M informative genes. Hence, the K_0 -dimensional embeddings of the n cells in the reference data are the columns of

$$\mathbf{S}_{(M)}^{\text{Ref}} = \mathbf{W}_{S,(M)}^T \times \mathbf{X}_{(M)} \in \mathbb{R}^{K_0 \times n}. \quad (6)$$

2. Denote the targeted gene profiling data of n' new cells with M informative genes measured by $\mathbf{X}_{(M)}^{\text{New}} \in \mathbb{R}_{\geq 0}^{M \times n'}$. Note that $\mathbf{X}_{(M)}^{\text{New}}$ contains log-transformed counts and has rows (genes) corresponding to the rows of $\mathbf{X}_{(M)}$. Project the n' cells to the K_0 -dimensional space by

$$\mathbf{S}_{(M)}^{\text{New}} = \mathbf{W}_{S,(M)}^T \times \mathbf{X}_{(M)}^{\text{New}} \in \mathbb{R}^{K_0 \times n'}. \quad (7)$$

3. (Optional) Normalize $\mathbf{S}_{(M)}^{\text{New}}$ and $\mathbf{S}_{(M)}^{\text{Ref}}$ to remove batch effects, if existent, by using a single-cell data integration method such as Harmony (Korsunsky et al., 2019).

Now the n reference cells and the n' new cells are in the same K_0 -dimensional space with biological variations preserved. Then, a classifier can be trained on the n reference cells' types and $\mathbf{S}_{(M)}^{\text{Ref}}$ for cell-type prediction, and it can be used to predict the n' cells' types from $\mathbf{S}_{(M)}^{\text{New}}$.

3 Results

3.1 scPNMF outputs a sparse and functionally interpretable representation of scRNA-seq data

We first demonstrate that scPNMF Step I, PNMF, outputs a sparse and functionally interpretable gene encoding of cells. We use the FregGold dataset (Freytag et al., 2018), which consists of three cell types (three human lung adenocarcinoma cell lines), and set the basis number $K=5$ for demonstration purpose. Both PCA and PNMF learn a weight matrix that can project the original scRNA-seq data

onto a 5D space. Unlike the weight matrix of PCA that has no zero entries, the weight matrix of PNMF is non-negative, highly sparse, containing 42.6% of entries as zeros, and has bases that are largely mutually exclusive (i.e. nonzero entries in different columns correspond to different rows/genes) (Fig. 2a). Compared with NMF, PNMF also has greater sparsity and mutual exclusiveness in bases (Supplementary Section S7). GO enrichment analysis shows that high weight genes in each PNMF basis are enriched with conceptually similar GO terms, and high weight genes in different PNMF bases are enriched with conceptually different GO terms (Fig. 2b). This result indicates that PNMF bases correspond to gene groups with distinct functions. On the contrary, the PCA bases do not have good functional interpretations: the high weight genes in each PCA basis are not enriched with conceptually similar GO terms, and different PCA bases share many high weight genes (Supplementary Fig. S8).

To further analyze the PNMF bases, we list the top 10 high weight genes in each basis (Supplementary Table S2), from which we identify many well-known genes with important functions. For instance, basis 1 contains classic housekeeping genes, such as *GAPDH* (Barber et al., 2005) and ribosomal protein genes (*RP*) (Silver et al., 2006); basis 3 contains well-known tumor-related genes, including *EGFR* (Blakely et al., 2017) and *CDK4* (O'Leary et al., 2016). In particular, the cells of the HCC827 cell line (one of the three cell types) have overall high scores in basis 3 (Supplementary Fig. S9), a reasonable result because the HCC827 cell line contains an *EGFR* activating mutation (Della Corte et al., 2017). In summary, scPNMF step I outputs bases representing sparse and functionally interpretable gene sets.

3.2 Basis selection is an essential step in scPNMF

Here, we explain why basis selection is an essential step in scPNMF. In the last section, we show that each PNMF basis of the FregGold dataset approximately represents one functional gene group. It is well known that housekeeping genes (basis 1) and cell-cycle genes (basis 4) are usually irrelevant to cell-type distinctions. However, such biological knowledge is not always available or certain. Therefore, scPNMF mainly relies on the two data-driven strategies: correlations with cell library sizes and multimodality tests (Section 2.2.2) for selecting informative bases.

Figure 2c visualizes the two strategies: cell scores in bases 1 and 4 are highly correlated with cell library sizes (Pearson correlations > 0.9); cell scores in bases 2 and 3 show strong evidence as multimodally distributed (adjusted P -value < 0.05). Hence, Strategy 1 will not retain bases 1 and 4, and Strategy 2 will not retain bases 1, 4 and 5; together, bases 1 and 4 will be removed, and bases 2, 3 and 5 will be selected. To verify the effectiveness of basis selection, we use UMAP to visualize cells based on the top 50 high weight genes in the unselected bases 1 and 4 versus those in the selected bases 2, 3 and 5 (Fig. 2d). We observe that the top genes in the unselected bases completely fail to separate the three cell types, while the top genes in the selected bases perfectly distinguish the three cell types. This result strongly supports that basis selection is a necessary step of scPNMF. If cell-type labels are provided, users may use a strategy alternative to 'correlations with cell library sizes' by regressing out the cell library sizes in a cell-type-specific manner from every basis (see Supplementary Section S6).

3.3 scPNMF outperforms state-of-the-art gene selection methods on diverse scRNA-seq datasets

In this section, we demonstrate scPNMF's capacity for informative gene selection. We comprehensively benchmark scPNMF against 11 other single-cell informative gene selection methods (Supplementary Table S3) on seven scRNA-seq datasets (Supplementary Table S4) using three clustering methods (Louvain clustering, K-means clustering and hierarchical clustering). For fair benchmarking, the seven scRNA-seq datasets cover both unique molecule identifier (UMI) and non-UMI protocols and include various biological samples. Using the adjusted Rand index (ARI) as the metric of clustering accuracy, we calculate the ARI values of the three clustering methods on each dataset using 100 informative genes selected by each gene

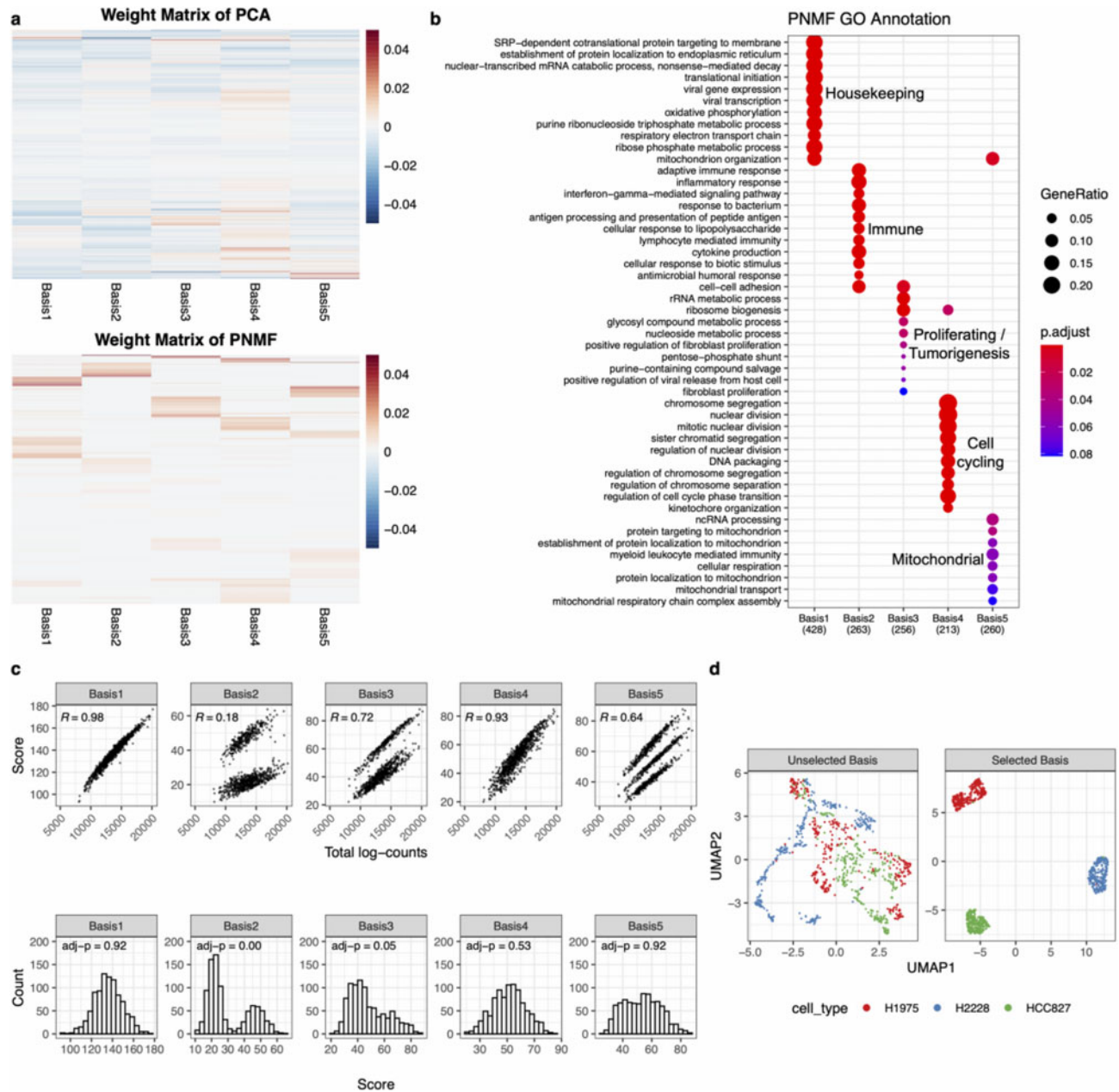


Fig. 2 Illustration of the sparse and interpretable projection found by scPNMF. We use the FregGold dataset as an example. (a) Comparison of the weight matrices of PCA and PNMF. Heatmaps visualize the learned weight matrices of PCA (top) and PNMF (bottom), where rows are genes and columns are bases. Red represents positive weights while blue represents negative weights. The rows are ordered by gene-wise hierarchical clustering. Compared to PCA, the weight matrix of PNMF is strictly non-negative, much more sparse and mutually exclusive between bases. (b) GO analysis result of each basis in the weight matrix of PNMF. Texts in black boxes summarize the functions of genes in each basis. The enriched GO terms are almost mutually exclusive, implying that each basis represents a unique gene functional cluster. (c) Statistical tests on each basis in the score matrix of PNMF. Top row: scatter plots of scores and total log-counts (cell library sizes). Each dot represents a cell. Cell scores in bases 1 and 4 are highly correlated with cell library sizes. Bottom row: histograms of cell scores in each basis. Scores in bases 2 and 3 show strong multimodality patterns (adjusted P -value ≤ 0.05). (d) UMAP visualizations of cells based on high weight genes in the unselected bases 1 and 4 and those in the selected bases 2, 3 and 5. Genes in the unselected bases completely fail to distinguish the three cell types, while genes in the selected bases lead to a clear separation of the three cell types.

selection method, as 100 genes are commonly used in targeted gene profiling.

Figure 3a shows that scPNMF has overall the highest ARI values across datasets and clustering methods. In particular, scPNMF has the highest average ARI value with each clustering method (Louvain: 0.83; K-means: 0.74; hierarchical clustering: 0.69) and the highest overall average ARI (0.75) across datasets and clustering methods. Note that the mean of the overall average ARI values of all methods except scPNMF is only 0.66.

We further show the UMAP visualization of cells in the Zheng4 dataset based on the informative genes selected by each of the 12

gene selection methods (Fig. 3b). Only scPNMF leads to a clear separation of naive cytotoxic T cells and regulatory T cells, while the informative genes selected by other methods except corFS and irlbaPcaFS cannot distinguish the two cell types at all.

We also compare the 12 methods under a varying number of informative genes: 20, 50, 200 and 500, the commonly used gene numbers in targeted gene profiling. We observe that the overall average ARI values of scPNMF are consistently higher than those of other methods, across all informative gene numbers (Supplementary Fig. S11). We apply the same benchmarking framework to scPNMF and its variant, where PNMF is replaced by NMF, and find that

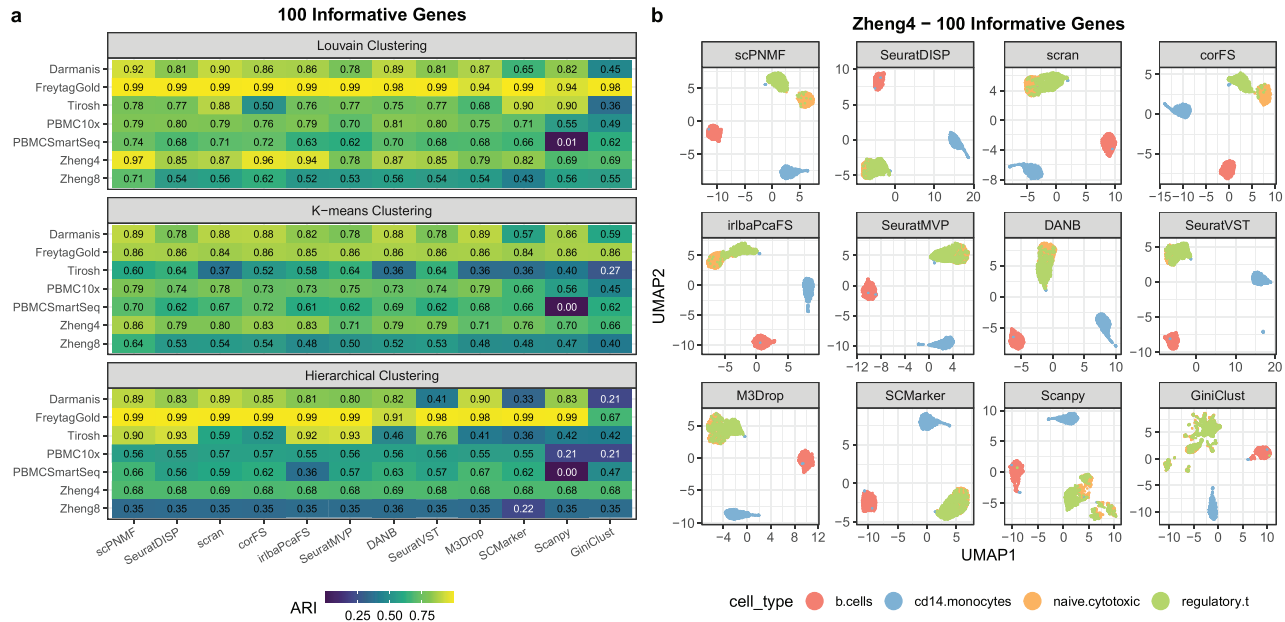


Fig. 3 Benchmarking scPNMF against 11 informative gene selection methods on seven scRNA-seq datasets. (a) Clustering accuracies (ARI values) of three clustering methods based on the informative genes selected. Gene selection methods are ordered from left to right by their average ARI across the three clustering methods and the seven datasets. (b) UMAP visualization of cells in the Zheng4 dataset based on 100 informative genes selected by each method. Genes selected by scPNMF lead to a clear separation between naive cytotoxic T cells and regulatory T cells, while the genes selected by others methods do not.

scPNMF performs consistently better (Supplementary Section S7). Moreover, compared with other methods, scPNMF leads to more stable overall average ARI values under varying numbers of informative genes, indicating its stronger robustness to the gene number constraint of targeted gene profiling. These results strongly support the superior performance of scPNMF as an informative gene selection method.

3.4 scPNMF guides targeted gene profiling experimental design and cell-type prediction

In this section, we demonstrate how scPNMF can guide the selection of genes to be measured in a targeted gene profiling experiment, and how scPNMF enables subsequent cell-type annotation on the targeted gene profiling data. We design two case studies with paired scRNA-seq reference data and ‘pseudo’ targeted gene profiling data, whose per-cell sequencing depth is higher than that of the corresponding scRNA-seq data.

In the first case study, we use the Zheng8 dataset (measured by the 10x Genomics protocol) as the reference dataset. To generate the pseudo targeted gene profiling data, we use a new single-cell gene expression simulator that captures gene correlations, scDesign2 (Sun et al., 2021), to generate data with a 100-time higher per-cell sequencing depth. In the second case study, we use the PBMC10x dataset (measured by 10x Genomics protocol) as the reference dataset, and we use PBMCSmartseq (measured by Smart-Seq2) as the pseudo targeted gene profiling data because Smart-Seq2 has a higher per-gene sequencing depth than 10x Genomics does. In both case studies, for each gene selection method, the corresponding pseudo targeted gene profiling datasets only contain the M informative genes selected by the method.

We benchmark scPNMF against the 11 gene selection methods in terms of cell-type prediction on the pseudo targeted gene profiling data. To avoid the bias for a specific classification algorithm, we apply three popular algorithms for cell-type prediction: random forest (Breiman, 2001), k-nearest neighbors (KNN) (Boser et al., 1992) and support vector machine (SVM) (Boser et al., 1992). In each case study, we first train each classification algorithm on the low-dimensional embeddings of the reference cells $S_{(M)}^{\text{Ref}}$ given the $M = 100$ informative genes selected by each gene selection method. Then, we

apply the trained classifier to the low-dimensional embeddings of the cells in the pseudo targeted gene profiling data $S_{(M)}^{\text{New}}$. Table 2 shows that scPNMF leads to the highest average prediction accuracy (0.81) across six combinations (two case studies $\times$ three classification algorithms). Moreover, scPNMF achieves the highest accuracy in each combination except Zheng8 + random forest where it is the second best. These results confirm that scPNMF effectively guides the selection of genes to measure in targeted gene profiling experiments, and it enables accurate cell-type annotation on newly generated targeted gene profiling datasets.

4 Discussion

We propose scPNMF, an unsupervised gene selection and data projection method for scRNA-seq data. The major goal of scPNMF is to select a fixed number of informative genes to distinguish cell types and guide gene selection for targeted gene profiling experiments. Moreover, scPNMF can project a new targeted gene profiling dataset with the selected genes to the low-dimensional space that embeds a reference scRNA-seq dataset. We perform a comprehensive benchmark to evaluate scPNMF in terms of informative gene selection against the state-of-the-art gene selection methods. Our results show that scPNMF consistently outperforms 11 existing methods for a wide range of informative gene numbers (from 20 to 500) on diverse scRNA-seq datasets. We also demonstrate that the informative genes selected by scPNMF can effectively guide gene selection for targeted gene profiling and lead to accurate cell-type annotation on targeted gene profiling data based on reference scRNA-seq data. In addition to the 11 methods, we compare scPNMF to the factorial single-cell latent variable model (f-sLVM) (Buettner et al., 2017), both conceptually and empirically, to clarify their differences and further illustrate the unique strength of scPNMF (see Supplementary Section S8).

Besides gene selection and data projection, scPNMF also works as a dimensionality reduction method with good interpretability. Each dimension in the low-dimensional space found by scPNMF can be considered as a new functional ‘feature’ (as a linear combination of correlated and thus functionally related genes). Moreover, the mutual exclusiveness makes the PNMF bases used in scPNMF advantageous over the PCA bases in terms of removing confounding

Table 2. Prediction accuracy of cell types based on 100 informative genes selected by 12 gene selection methods in the two case studies with paired reference scRNA-seq data and targeted gene profiling data

Method	Zheng8			PBMC			Average Accuracy
	RandomForest	KNN	SVM	RandomForest	KNN	SVM	
scPNMF	0.85 (0.83,0.87)	<u>0.80</u> (0.78,0.83)	<u>0.87</u> (0.85,0.89)	<u>0.84</u> (0.79,0.88)	<u>0.84</u> (0.79,0.88)	<u>0.67</u> (0.61,0.73)	<u>0.81</u>
M3Drop	0.85 (0.83,0.87)	<u>0.80</u> (0.77,0.83)	<u>0.87</u> (0.84,0.89)	<u>0.84</u> (0.79,0.88)	0.77 (0.71,0.82)	0.63 (0.57,0.69)	0.79
SeuratDISP	0.84 (0.81,0.86)	0.78 (0.75,0.81)	0.86 (0.84,0.88)	0.80 (0.75,0.84)	0.75 (0.70,0.80)	0.64 (0.58,0.70)	0.78
corFS	0.80 (0.77,0.82)	0.75 (0.73,0.78)	0.82 (0.80,0.85)	0.82 (0.77,0.86)	0.81 (0.76,0.86)	0.62 (0.56,0.68)	0.77
GiniClust	<u>0.86</u> (0.83,0.88)	0.79 (0.76,0.81)	0.86 (0.83,0.88)	0.80 (0.75,0.84)	0.76 (0.71,0.81)	0.53 (0.47,0.60)	0.75
scan	0.79 (0.76,0.81)	0.72 (0.69,0.75)	0.82 (0.80,0.85)	0.78 (0.72,0.82)	0.73 (0.67,0.78)	0.67 (0.61,0.72)	0.75
SeuratMVP	0.83 (0.81,0.85)	0.77 (0.74,0.80)	0.85 (0.82,0.87)	0.82 (0.77,0.86)	0.74 (0.69,0.79)	0.47 (0.40,0.53)	0.74
Scanpy	0.79 (0.77,0.82)	0.71 (0.68,0.74)	0.80 (0.78,0.83)	0.80 (0.75,0.84)	0.76 (0.71,0.81)	0.52 (0.46,0.58)	0.73
SCMarker	0.77 (0.74,0.79)	0.68 (0.65,0.71)	0.74 (0.71,0.77)	0.77 (0.71,0.81)	0.71 (0.65,0.76)	0.45 (0.39,0.52)	0.69
SeuratVST	0.73 (0.70,0.76)	0.68 (0.65,0.71)	0.75 (0.73,0.78)	0.74 (0.68,0.79)	0.68 (0.63,0.74)	0.40 (0.34,0.46)	0.67
DANB	0.71 (0.68,0.73)	0.69 (0.66,0.71)	0.75 (0.73,0.78)	0.73 (0.67,0.78)	0.74 (0.68,0.79)	0.28 (0.23,0.34)	0.65
irlbaPcaFS	0.68 (0.65,0.71)	0.61 (0.58,0.64)	0.71 (0.68,0.74)	0.71 (0.65,0.76)	0.77 (0.71,0.82)	0.16 (0.12,0.21)	0.61

Parentheses are 95% confidence intervals. Highest number within each column is labeled by boldface and underline.

effects. For example, cell-cycle genes obscure the identification of cell types and should be removed from low-dimensional embeddings of cells. For PCA, cell-cycle genes affect many PCA bases, so the popular scRNA-seq pipeline Seurat implements a complicated approach that first calculates ‘cell-cycle scores’ and then regresses each basis (principal component) on these scores to remove the effects of cell-cycle genes (Stuart *et al.*, 2019). In contrast, cell-cycle genes are concentrated in only one PNMF basis, so it is easy to remove that basis to clear the effects of cell-cycle genes. Therefore, scPNMF has great potentials in deciphering cell heterogeneity in single-cell data by working as an interpretable dimensionality reduction method.

The current implementation of scPNMF focuses on single-cell gene expression data. Considering the rapid development of single-cell multiomics technologies, we plan to extend scPNMF to accommodate other technologies that measure other genomics features such chromatin accessibility landscapes measured by single-cell ATAC-seq (Pott and Lieb, 2015), or even to integrate data across multiomics datasets. Another note is that the multimodality test for basis selection in scPNMF only accounts for discrete cell types but not continuous cell trajectories. Therefore, other tests or strategies are needed to select informative bases to capture biological variations along continuous cell trajectories.

An important question for gene selection is: how many genes should be selected as informative genes to fully capture the biological variations of interest? In our studies, we observe that, after the informative gene number reaches 200, the clustering accuracies based on the selected informative genes plateau for most gene selection methods including scPNMF. Therefore, 200 genes may be sufficient for capturing biological variations in scRNA-seq data. However, it remains challenging to decide the minimum number of informative genes, given that the underlying cell subpopulation structure is data-specific and might be complex. We plan to explore this problem in the future with the possible use of information theory.

Financial Support: This work was supported by the following grants: National Science Foundation DBI-1846216, NIH/NIGMS R01GM120507, Johnson & Johnson WiSTEM2D Award, Sloan Research Fellowship, and UCLA David Geffen School of Medicine W.M. Keck Foundation Junior Faculty Award (to J.J.L); NIH/NINDS R01NS117148 (to R.W).

Conflict of Interest: none declared.

References

Ameijeiras-Alonso, J. *et al.* (2019) Mode testing, critical bandwidth and excess mass. *Test*, **28**, 900–919.

Andrews, T.S., and Hemberg, M. (2019) M3drop: dropout-based feature selection for scRNAseq. *Bioinformatics*, **35**, 2865–2867.

Barber, R.D. *et al.* (2005) GAPDH as a housekeeping gene: analysis of GAPDH mRNA expression in a panel of 72 human tissues. *Physiol. Genomics*, **21**, 389–395.

Baron, M. *et al.* (2016) A single-cell transcriptomic map of the human and mouse pancreas reveals inter- and intra-cell population structure. *Cell Syst.*, **3**, 346–360.

Birnbaum, K.D. (2018) Power in numbers: single-cell RNA-seq strategies to dissect complex tissues. *Annu. Rev. Genetics*, **52**, 203–221.

Blakely, C.M. *et al.* (2017) Evolution and clinical impact of co-occurring genetic alterations in advanced-stage egfr-mutant lung cancers. *Nat. Genetics*, **49**, 1693–1704.

Boileau, P. *et al.* (2020) Exploring high-dimensional biological data with sparse contrastive principal component analysis. *Bioinformatics*, **36**, 3422–3430.

Boser, B.E. *et al.* (1992) A training algorithm for optimal margin classifiers. In: *Proceedings of the Fifth Annual Workshop on Computational Learning Theory (COLT ’92)*. Association for Computing Machinery, New York, NY, USA, 144–152.

Breiman, L. (2001) Random forests. *Mach. Learn.*, **45**, 5–32.

Brunet, J.-P. *et al.* (2004) Metagenes and molecular pattern discovery using matrix factorization. *Proc. Natl. Acad. Sci. USA*, **101**, 4164–4169.

Buettner, F. *et al.* (2017) f-scLVM: scalable and versatile factor analysis for single-cell RNA-seq. *Genome Biol.*, **18**, 1–13.

Della Corte, C.M. *et al.* (2017) Efficacy of continuous EGFR-inhibition and role of hedgehog in egfr acquired resistance in human lung cancer cells with activating mutation of EGFR. *Oncotarget*, **8**, 23020–23032.

Duren, Z. *et al.* (2018) Integrative analysis of single-cell genomics data by coupled nonnegative matrix factorizations. *Proc. Natl. Acad. Sci. USA*, **115**, 7723–7728.

Durif, G. *et al.* (2019) Probabilistic count matrix factorization for single cell expression data analysis. *Bioinformatics*, **35**, 4011–4019.

Eisenberg, E., and Levanon, E.Y. (2013) Human housekeeping genes, revisited. *Trends Genetics*, **29**, 569–574.

Freytag, S. *et al.* (2018) Comparison of clustering tools in R for medium-sized 10x genomics single-cell RNA-sequencing data. *F1000Research*, **7**, 1297.

Gao, C., and Welch, J.D. (2020). Iterative refinement of cellular identity from single-cell data using online learning. In: *International Conference on Research in Computational Molecular Biology*. p. 248–250. Springer, Cham.

Hafemeister, C., and Satija, R. (2019) Normalization and variance stabilization of single-cell rna-seq data using regularized negative binomial regression. *Genome Biol.*, **20**, 1–15.

Jiang, L. *et al.* (2016) GiniClust: detecting rare cell types from single-cell gene expression data with gini index. *Genome Biol.*, **17**, 144.

Korsunsky, I. *et al.* (2019) Fast, sensitive and accurate integration of single-cell data with harmony. *Nat. Methods*, **16**, 1289–1296.

Lee, D.D., and Seung, H.S. (1999) Learning the parts of objects by non-negative matrix factorization. *Nature*, **401**, 788–791.

- Lun,A.T. *et al.* (2016) Pooling across cells to normalize single-cell RNA sequencing data with many zero counts. *Genome Biol.*, **17**, 75.
- Macosko,E.Z. *et al.* (2015) Highly parallel genome-wide expression profiling of individual cells using nanoliter droplets. *Cell*, **161**, 1202–1214.
- Marshall,J.L. *et al.* (2020) HyPR-seq: single-cell quantification of chosen RNAs via hybridization and sequencing of DNA probes. *Proc. Natl. Acad. Sci. USA*, **117**, 33404–33413.
- Moffitt,J.R. *et al.* (2016) High-throughput single-cell gene-expression profiling with multiplexed error-robust fluorescence in situ hybridization. *Proc. Natl. Acad. Sci. USA*, **113**, 11046–11051.
- O’Leary,B. *et al.* (2016) Treating cancer with selective cdk4/6 inhibitors. *Nat. Rev. Clin. Oncol.*, **13**, 417–430.
- Pott,S., and Lieb,J.D. (2015) Single-cell ATAC-seq: strength in numbers. *Genome Biol.*, **16**, 1–4.
- Potter,S.S. (2018) Single-cell RNA sequencing for the study of development, physiology and disease. *Nat. Rev. Nephrol.*, **14**, 479–492.
- Raj,A. *et al.* (2008) Imaging individual mrna molecules using multiple singly labeled probes. *Nat. Methods*, **5**, 877–879.
- Silver,N. *et al.* (2006) Selection of housekeeping genes for gene expression studies in human reticulocytes using real-time PCR. *BMC Mol. Biol.*, **7**, 33.
- Stuart,T. *et al.* (2019) Comprehensive integration of single-cell data. *Cell*, **177**, 1888–1902.
- Subramanian,A. *et al.* (2017) A next generation connectivity map: L1000 platform and the first 1,000,000 profiles. *Cell*, **171**, 1437–1452.
- Sun,T. *et al.* (2021). scDesign2: a transparent simulator that generates high-fidelity single-cell gene expression count data with gene correlations captured. *Genome Biol.*, **22**, 163.
- Thellin,O. *et al.* (1999) Housekeeping genes as internal standards: use and limits. *J. Biotechnol.*, **75**, 291–295.
- Uzbas,F. *et al.* (2019) Bart-seq: cost-effective massively parallelized targeted sequencing for genomics, transcriptomics, and single-cell analysis. *Genome Biol.*, **20**, 1–16.
- Wang,F. *et al.* (2019) Scmarker: ab initio marker selection for single cell transcriptome profiling. *PLoS Comput. Biol.*, **15**, e1007445.
- Welch,J.D. *et al.* (2019) Single-cell multi-omic integration compares and contrasts features of brain cell identity. *Cell*, **177**, 1873–1887.
- Yang,Z., and Michailidis,G. (2016) A non-negative matrix factorization method for detecting modules in heterogeneous omics multi-modal data. *Bioinformatics*, **32**, 1–8.
- Yang,Z., and Oja,E. (2010) Linear and nonlinear projective nonnegative matrix factorization. *IEEE Trans. Neural Netw.*, **21**, 734–749.
- Ye, Y., and Li, J.J. (2016) NMFP: a non-negative matrix factorization based preselection method to increase accuracy of identifying mRNA isoforms from RNA-seq data. *BMC Genomics.*, **17**, 11. <https://doi.org/10.1186/s12864-015-2304-8>
- Yuan,Z. *et al.* (2009) Projective nonnegative matrix factorization: sparseness, orthogonality, and clustering. *Neural Process. Lett.*, **11**–13. [CVOCROSSCVO]
- Zhang,S. *et al.* (2020) Dimensionality reduction for single cell RNA sequencing data using constrained robust non-negative matrix factorization. *NAR Genomics Bioinformatics*, **2**, lqaa064.
- Zhu,C. *et al.* (2020) Single-cell multimodal omics: the power of many. *Nat. Methods*, **17**, 11–14.
- Zhu,X. *et al.* (2017) Detecting heterogeneity in single-cell RNA-seq data by non-negative matrix factorization. *PeerJ*, **5**, e2888.