

A unified framework for selecting and evaluating cell-type-specific gene co-expressions in single-cell data

Xinning Shan , Yingxin Lin, Hongyu Zhao *

Department of Biostatistics, Yale University, New Haven, CT 06511, United States

*Corresponding author. E-mail: hongyu.zhao@yale.edu

Abstract

Cell-type-specific gene co-expression networks are widely used to characterize gene relationships. Although many methods have been developed to infer such co-expression networks from single-cell data, the lack of consideration of false positive control in many evaluations and downstream analyses may lead to incorrect conclusions because higher reproducibility, higher functional coherence, and a larger overlap with known biological networks may not imply better performance if the false positives are not well controlled. In this study, we systematically compared two distinct criteria for selecting correlated gene pairs from single-cell data, p -value versus correlation strength. We found that the use of p -values instead of correlation strength is more robust for both selecting meaningful gene pairs and for the fair benchmarking of co-expression estimation methods. To make this approach universally applicable, we extended and validated a simulation method that can efficiently and reliably generate empirical p -values for co-expression estimation methods that do not have corresponding or well-controlled p -values. Furthermore, we demonstrated that a fair comparison of the estimation methods requires adjusting for the varying number of gene pairs they identified and accounting for the inherent expression-level biases within ground truth biological networks. Our study provides a practical guide for researchers to select reliable correlated gene pairs for downstream study and establishes a more rigorous standard for the evaluation and comparison of gene co-expression network estimation methods.

Keywords single cell, cell-type-specific, co-expression networks

Background

With the development of single-cell technology and the potential of inferring cell-type-specific gene co-expression networks from these data to understand relationships among genes, many methods have been developed for estimating such networks through single-cell RNA-sequencing (scRNA-seq) data. These methods typically result in a dense network with a value for each possible gene pair. Therefore, the accurate identification of co-expressed gene pairs from the initial network is a critical step for downstream analyses and method benchmarking. However, this post-selection step introduces two key challenges. First, there is no consensus approach for selecting truly co-expressed gene pairs. Second, there is no standard for the fair comparison of the gene co-expression estimation methods.

There are two commonly used approaches for selecting correlated gene pairs in real data analysis, p -value-based and correlation-strength-based. The correlation-strength-based approach selects correlated gene pairs based on the estimated correlation strength, and is more commonly used because of its convenience [1–8]. However, the correlation-strength-based approach does not consider false positive control. Although a p -value-based approach is theoretically

more robust, a formal comparison to quantify the practical difference between these two approaches is lacking in the context of cell-type-specific single-cell analysis to the best of our knowledge. Furthermore, a rigorous p -value-based approach is hampered by the fact that many methods do not provide p -values (e.g. propr [9]) or the p -value is not well-controlled (e.g. Pearson/Spearman correlation on log-normalized data, Noise Regularization [6], and Pearson correlation on the analytic Pearson residuals [10]) as shown in Su *et al.* [11]. These issues are further complicated during method evaluation. Standard benchmarks are often biased not only because different methods identify widely varying numbers of correlated gene pairs, but also because the biological databases used as ground truth can have their own expression biases, making direct comparisons unfair.

To address these critical gaps, we propose a unified, p -value-driven framework for post-selecting and evaluating cell-type-specific gene co-expressions. We first establish our primary finding that p -value-based thresholding is superior to correlation-strength-based thresholding using one exemplar of the co-expression network estimation method, which uniquely provides both estimates and well-calibrated p -values (Section 2.1). To generalize this finding, we address the lack of

Received: December 9, 2025. **Revised:** April 18, 2026. **Accepted:** May 12, 2026

© The Author(s) 2026. Published by Oxford University Press.

This is an Open Access article distributed under the terms of the Creative Commons Attribution-NonCommercial License (<https://creativecommons.org/licenses/by-nc/4.0/>), which permits non-commercial re-use, distribution, and reproduction in any medium, provided the original work is properly cited. For commercial re-use, please contact reprints@oup.com for reprints and translation rights for reprints. All other permissions can be obtained through our RightsLink service via the Permissions link on the article page on our site—for further information please contact journals.permissions@oup.com.

well-calibrated p -values in other co-expression estimation methods by generating empirical p -values. To do this robustly, we first conduct a benchmark of simulation methods and identify scSimu as the optimal choice for this task (Section 2.2). Next, using scSimu to generate well-calibrated empirical p -values for all methods, we generalize our initial finding in Section 2.1 (Section 2.3). Finally, we extend this framework to enable fair method comparison, introducing a necessary adjustment that accounts for method-specific differences in statistical power (Section 2.4). Our work provides both a practical workflow for researchers and a more rigorous standard for the benchmarking of gene co-expression estimation methods.

Results

A lack of type I error control in identifying correlated gene pairs can lead to misleading results

In this section, we demonstrate the critical need for using p -values to control false positives in the evaluation of gene co-expression network estimation methods. We focused our analysis on CS-CORE, a method that provides well-calibrated p -values for cell-type-specific gene co-expressions [11]. Because we do not know the ground truth of correlated gene pairs in real data, we used simulations that mimicked real data to better understand the impact of appropriate control of false positives (see Methods for details). For the evaluation measures, we focused on two most commonly used ones, reproducibility [1, 4, 5, 12] and overlap with known biological networks (STRING [13] and Reactome [14]) [1, 2, 4, 6, 7, 15].

We used these simulations to directly compare the performance of the correlation-strength and p -value-based approaches. In the correlation-strength-based approach, selecting a larger number of top pairs leads to higher reproducibility and more overlaps with biological networks (Figs 1A and 1F). However, it may include more false positives. As shown in Figs 1A and 1F, an increasing number of reproducible pairs and a higher number of overlaps with STRING are both associated with a lower precision. We also observe higher reproducibility and more overlaps when the p -value-based approach with a less stringent cut-off is used (Figs 1B and 1G), but there is a less decline in precision (Figs 1B and 1G).

To understand the influence of these false positives in the evaluation process, we investigated the proportion of misidentified reproducible pairs (see Methods for details). As expected, there is an increased proportion of misidentification with a growing number of gene pairs selected (Fig. 1C). We find that the proportion of misidentified reproducible pairs based on p -value-based approach is much lower than that based on the correlation-strength-based approach (Fig. 1D). The results are similar to those comparing overlaps with STRING (Figs 1H–I) and overlaps with Reactome (Fig. S2). Even though the correlation-strength-based approach has a much higher number of reproducible pairs and overlaps with known biological networks, the large discrepancy in precision and proportion of misidentification at different thresholds necessitates more careful threshold selection. Additionally, variations in the density and magnitude of the correlation matrix across datasets further complicate threshold selection, making applying the same threshold to different datasets questionable.

As shown in Fig. 1, there is a potential trade-off between the proportion of false identification and the number of true reproducible pairs or true overlaps with known biological networks, i.e. a lower

proportion of false identification may correspond to a lower number of true reproducible pairs or true overlaps. To better quantify this, we compared the number of true reproducible pairs identified by the correlation-strength-based approach and the p -value-based approach, ensuring that their proportions of false identifications are the same in Fig. 1E. This comparison is functionally equivalent to a precision-recall analysis and further supports that selecting correlated gene pairs based on p -value is better than that based on correlation strength. Similarly, we observe a higher number of true overlaps with known biological networks when using the p -value-based approach (Figs 1J and S2E).

A validated simulation to enable p -value selection across all methods

While our results demonstrated the superiority of using p -value-based approach, its practical application is often limited as many co-expression estimation methods do not provide p -values. Among those that do, p -values are often derived from theoretical null distributions, such as the Student's t -distribution, which may not be suitable for sparse and complex scRNA-seq data. This is further supported by Su *et al.* [11], which examined the p -value distribution of many co-expression estimation methods on permuted single-cell RNA count data. It was found that only Normalizr [16], Pearson correlation on the sctransform normalized data [17], and CS-CORE have well-calibrated type I errors [11].

Even though Pearson correlation on sctransform normalized data was shown to have well-calibrated type I errors, the p -values are derived from a theoretical Student's t -distribution. As shown by $t = \frac{r\sqrt{n-2}}{\sqrt{1-r^2}}$, this p -value only depends on the correlation coefficient (r) and the number of single cells used (n). This means that two gene pairs with the same r and n will yield identical p -values, regardless of other inherent differences in their expression distributions, such as mean and variance (Figs S3A and S3B). There are similar concerns for the Spearman correlation. Therefore, a data-driven empirical p -value, which can account for these specific data characteristics, is necessary to provide a more robust and appropriate measure of statistical significance. To address this need, we employed a simulation-based framework to generate well-controlled empirical p -values, thereby enabling a rigorous and significance-based selection for any co-expression estimation methods.

Many methods have been developed for simulating count data in the context of scRNA-seq. Crowell *et al.* [18] found that ZINB-WaVE [19] and scDesign2 [20] were the top two methods for simulating data containing cells from a single batch and cluster. In another benchmark paper, ZINB-WaVE and SPARSim [21] were found to exhibit greater similarity to the original data in terms of both data properties and biological signals [22]. More recently, scDesign2 has been further developed into scDesign3 [23]. Even though these methods can closely mimic real data, some limitations exist. ZINB-WaVE [19] does not consider the correlation structure, whereas scDesign2 [20] and scDesign3 [23] are computationally demanding requiring several hours for a single simulation (Figs 2C and 2D). SPARSim's [21] assumption of a fixed sequencing depth per cell is a key limitation, as this does not align with the stochastic nature of sequencing in real datasets. Therefore, we introduce scSimu (Single-Cell SIMulation), an improved and extended version of the simulation method used in our previous publication [11] (see Methods for details), which serves as an accurate

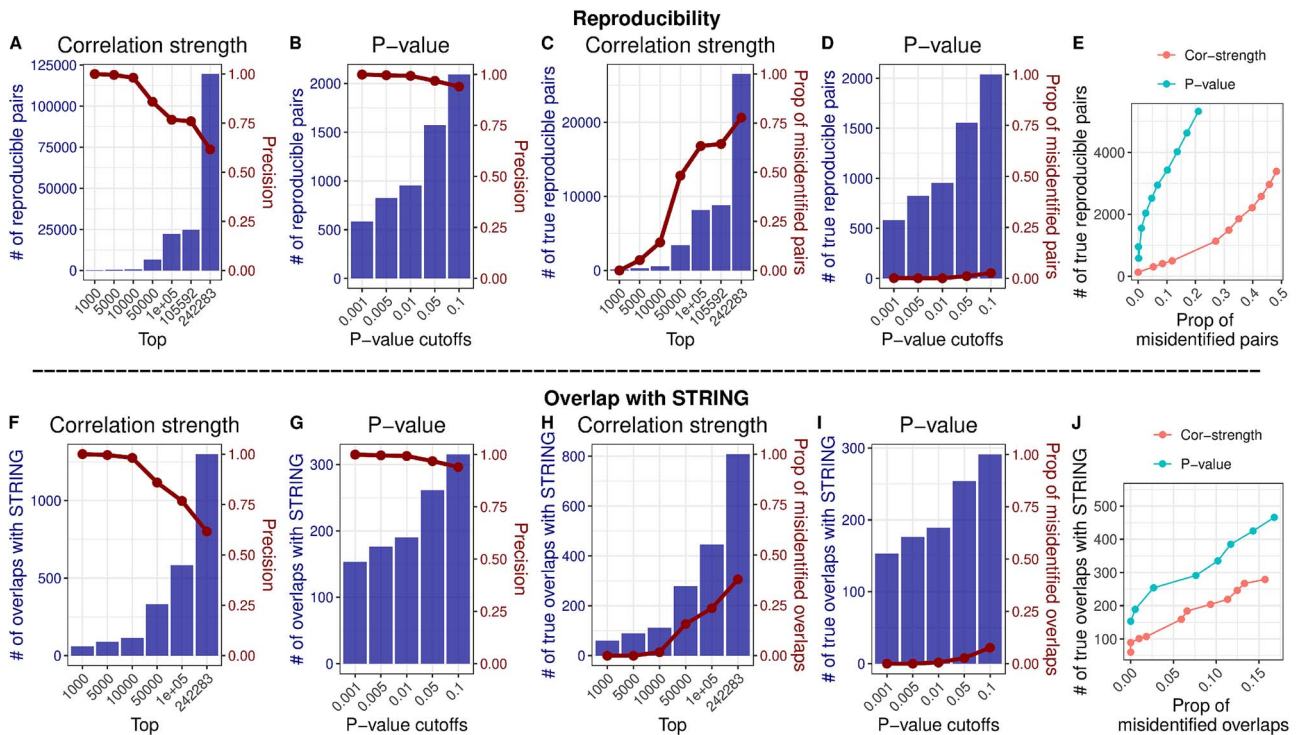


Figure 1 Comparison of correlation-strength and p -value-based approaches for selecting co-expressed gene pairs. Using results from CS-CORE, the performance of the correlation-strength-based and the p -value-based approach was evaluated by reproducibility (A–E) and overlap with the STRING database (F–J). A–B. Total number of reproducible pairs (blue bars, left y-axis) and precision (red line, right y-axis) for the correlation-strength (A) and p -value (B) approaches. The x-axis represents the number of top pairs selected by correlation strength or the p -value cutoff. (Precision in PNAS data can be found in Fig. S1) C–D. Number of true reproducible pairs (blue bars, left y-axis) and the proportion of misidentified pairs, $1 - \frac{\text{# of true reproducible pairs}}{\text{# of reproducible pairs}}$ (red line, right y-axis). E. Direct comparison of the number of true reproducible pairs versus the proportion of misidentified pairs. F–G. Total number of overlaps with the STRING and precision for the correlation-strength (F) and p -value (G) approaches. H–I. Number of true overlaps with STRING and the proportion of misidentified overlaps for the correlation-strength (H) and p -value (I) approaches. J. Direct comparison of the number of true overlaps versus the proportion of misidentified overlaps.

and efficient tool for evaluating co-expression networks by simulating realistic scRNA-seq count data under the null hypothesis.

To validate the accuracy of scSimu in mimicking real data, we benchmarked its ability to simulate both independent and correlated genes against a total of 13 other simulation methods (Table S1; see Methods for details). As a baseline check, we first confirmed that all simulation methods in the independent comparison successfully generated a null dataset of independent genes (Table S2 and Fig. S4). As shown in Figs 2A and 2B, scSimu, SPARSim, and scDesign3 achieve the greatest similarity with real data for both tasks. This is in line with the observation of Song *et al.* [23] that scDesign3 had better performance and that of Cao *et al.* [22] where SPARSim was the second best method. However, ZINB-WaVE, the best method in two benchmark papers [18, 22], is not among the top three methods in our comparisons. Our benchmark exclusively evaluates methods on data from droplet-based technologies, as their high throughput and cost-effectiveness make them widely used in the field. The model assumption for scRNA-seq data differs across sequencing platforms. For example, a zero-inflated model is used for the Smart-Seq data [24, 25]. In the two benchmark papers [18, 22], different types of scRNA-seq data were considered, likely contributing to conclusions that were inconsistent with ours. Although scSimu, SPARSim, and scDesign3 show similar performance, SPARSim's generating mechanism is unrealistic and scDesign3's extensive computational time makes it impractical for simulating null distributions (Figs 2C and 2D). Separated results for the KDE and KS tests (Fig. S5) confirm these findings.

Based on its balance of accuracy and computational efficiency, we selected scSimu as the null data generator for our p -value framework. To provide a further validation for this purpose, we examined the distribution of empirical p -values of independent gene pairs from a scSimu generated real data simulation. As shown in Fig. 3, the resulting p -value distribution for independent gene pairs almost aligns with a Uniform(0,1) distribution. Furthermore, this holds true for all seven gene co-expression estimation methods we considered, regardless of whether they originally provided p -values. This demonstrates that scSimu provides effective and consistent control over type I error, making it a reliable tool for generating the empirical p -values needed to test for significant co-expression.

p -value-based approach consistently outperforms and avoids benchmark distortions

To determine if the superiority of the p -value approach is a general principle, we extended our comparison to several other co-expression estimation methods. In Section 2.1, we showed that a p -value-based approach is better than the correlation-strength-based approach under CS-CORE. In this section, we extended this argument to a more general case by comparing the correlation-strength-based approach and the p -value-based approach using empirical p -values for six additional methods, including propr, Noise Regularization, Pearson correlation on log-transformed data (Pearson), Spearman correlation on log-transformed data (Spearman), Pearson correlation on the

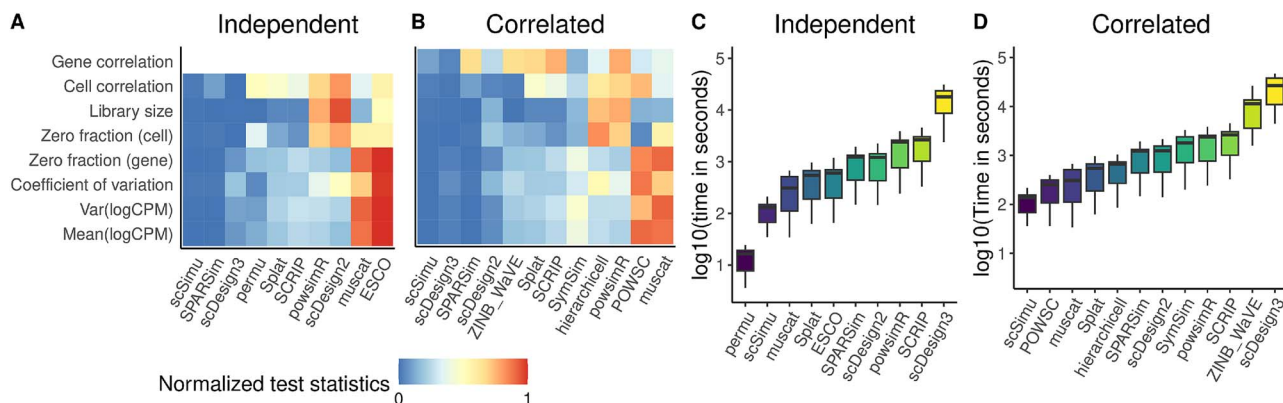


Figure 2 Benchmark of the simulation methods. A and B. The x-axis is the simulation method. The y-axis is the evaluation metric. The color represents the normalized test statistics of the combination of the KDE test and the KS test. Methods on the x-axis are ordered by the average normalized test statistics. A. Performance of simulating independent genes. B. Performance of simulating correlated genes. C and D. Boxplots of the computational time for simulating independent and correlated genes respectively, across 10 datasets. The y-axis represents the computational time, measured in seconds and displayed on a log10 scale. We specified the marginal distribution as negative binomial to improve the speed in scDesign2.

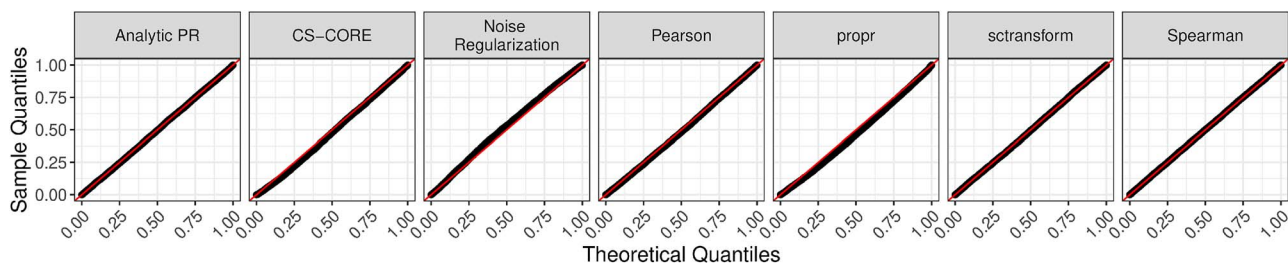


Figure 3 Q-Q plots of empirical p -values for independent gene pairs. Empirical p -values were calculated according to the methods outlined in Section 4.4.

analytic Pearson residuals (Analytic PR), and Pearson correlation on the normalized data obtained via sctransform (sctransform). We also included CS-CORE using empirical p -values to further test if the previously drawn conclusion is specific to CS-CORE provided p -values.

The results are consistent with the findings from Section 2.1. For every method tested, the correlation-strength-based approach identifies a larger number of reproducible pairs and overlaps with biological networks, but at the cost of significantly lower precision and a higher proportion of misidentified pairs (Figs 4, 5, 6, S7, S8, and S9). In addition, the proportion of misidentification is especially high for certain gene co-expression estimation methods (Fig. 4D). This finding indicates that some gene co-expression estimation methods, when using a correlation-strength-based approach, may seem artificially more effective due to not accounting for false positives. We estimated empirical p -values through scSimu and evaluated the p -value-based approach. Conversely, the p -value-based approach consistently maintains a lower proportion of false identifications, even when the p -value cutoff is relaxed to select more pairs (Figs 4, 5, 6, S7, S8, and S9). We further validated these results using a cell-type-specific TF-target database and observed similar phenomena (S10 and S11). This confirms that the advantage of using statistical significance to control for false positives is a robust and generalizable finding, not specific to a single co-expression estimation method.

Without considering the p -value in selecting correlated gene pairs can also affect the benchmark results across gene co-expression estimation methods. The ranks of estimation methods at the same

cutoff vary substantially in the number of reproducible pairs and the number of true reproducible pairs when we use the correlation-strength-based approach (Figs 4A and 4C), i.e. a method that exhibits the highest reproducibility may not be the best method since many of the reproducible pairs may be false positives. On the contrary, the ranks are rather consistent when we use the p -value-based approach (Figs 4E and 4G). We observe the same pattern through overlaps with known biological networks (Figs 5 and S7). Our observation also explains the inconsistent results between the benchmark study of gene association estimation methods conducted by Skinnider *et al.* [1], which found propr to be superior, and the study conducted by Su *et al.* [11] where propr did not perform well, due to the lack of controls of false positives. Our study highlights the importance of considering statistical significance in gene pair selection. By selecting correlated gene pairs based on the p -value, as we demonstrated above, more reliable evaluation results can be achieved.

We further compared the two approaches for selecting correlated gene pairs using real data, focusing on oligodendrocytes and excitatory neurons in the ROSMAP data. To perform the comparison, we calculated the number of shared Gene Ontology (GO) terms for gene pairs identified exclusively by either the p -value-based approach or the correlation-strength-based approach. A higher number of shared GO terms indicates a stronger likelihood that two genes are correlated. Overall, gene pairs uniquely identified by the p -value-based approach exhibit a higher number of shared GO terms, suggesting that the p -value-based approach identifies more reliable correlated gene pairs

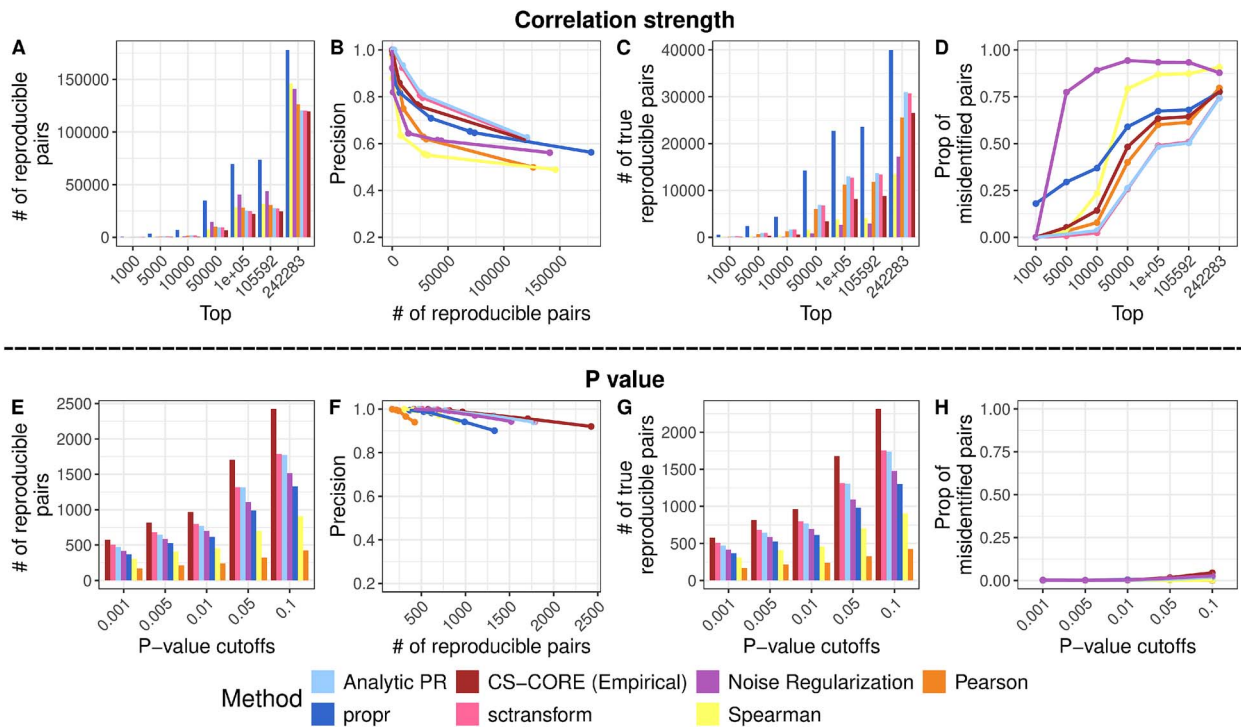


Figure 4 Performance in terms of reproducibility. Different estimation methods are represented with different colors. We use the term CS-CORE (Empirical) to denote the CS-CORE when its p -values is from our scSimu framework. Correlated gene pairs were identified via the correlation-strength-based approach (A–D), and p -value-based approach (E–H). A and E. The number of identified reproducible pairs at different cutoffs. B and F. Relationship between precision in ROSMAP data (precision in PNAS data can be found in Fig. S6) and the number of reproducible pairs. Points in B and F correspond to different cutoffs for the correlation strength and p -value, respectively. C and G. The number of true reproducible pairs at different cutoffs. D and H. The proportion of misidentified reproducible pairs when gene pairs were selected by their estimated correlations (D) or by the empirical p -values (H).

(Figs 7 and S12). Furthermore, through clustering analysis of the identified correlated gene pairs in excitatory neurons, we found that the p -value-based approach identifies gene clusters more directly related to the excitatory neuron functions (Fig. S13).

Consideration of the overlap by chance and bias in the biological networks in the evaluation process

The p -value-based approach will cause a large discrepancy in the number of identified correlated gene pairs between different estimation methods. This introduces a critical bias in performance evaluations, as methods that identify more pairs will have a higher chance of random overlap with known biological databases or results from other replicate datasets. Therefore, a fair comparison necessitates an adjusted framework that normalizes for these differences to account for this statistical artifact.

In addition to the statistical artifact of random overlaps, a fair evaluation is further complicated by a previously overlooked issue, the inherent expression bias within the biological databases that commonly used for validation. We observe that the known biological networks are biased toward genes and gene pairs with certain expression levels. To illustrate this, we calculated the average expression levels of the genes based on the lung dataset [26]. We contrast the expression levels of genes within biological networks versus the other genes and find that genes in biological networks are biased toward a higher expression level (Fig. 8A and 8D). We further consider pair-level bias. We established the background expression level for gene pairs by calculating the geometric mean of the expression levels for every

possible combination, using the genes present in the biological network. As shown in Figs 8B and 8E, the genes in the biological network, which are already biased in expression levels, exhibit an additional tendency toward higher expression at the pair level. Although there is a general tendency for genes and gene pairs in known biological networks to have higher expression levels, the degree of such bias varies across datasets. For example, while there is an expression bias across a broad range of genes for the lung dataset (Fig. S14), the bias is negligible for the top 1000 or 5000 genes in the ROSMAP dataset (Figs S15 and S16). This dataset-specific nature of the bias means that a one-size-fits-all correction is inadequate and necessitates a more flexible approach.

To address the biases from both random overlaps and expression levels, we developed a flexible adjustment strategy based on a stratified hypergeometric model. This model allows us to estimate the expected number of random overlaps and its standard deviation after accounting for expression-level bias, ultimately providing a standardized Z-score for evaluation. We first categorize gene pairs into k expression bins based on their mean expression levels. Within any given bin i , the number of overlaps that would occur by chance can be modeled precisely. Let n_i be the total number of gene pairs in bin i , b_i be the number of those pairs that are in the known biological network, and c_i be the number of identified correlated gene pairs. Then, the number of random overlap with the biological network in bin i follows a hypergeometric distribution, $\text{Hypergeometric}(n_i, b_i, c_i)$. The total expected random overlap and its variance can be calculated by summing the expectations and variances from all k bins as the sampling process is independent across bins. From this, we derive a Z-score that

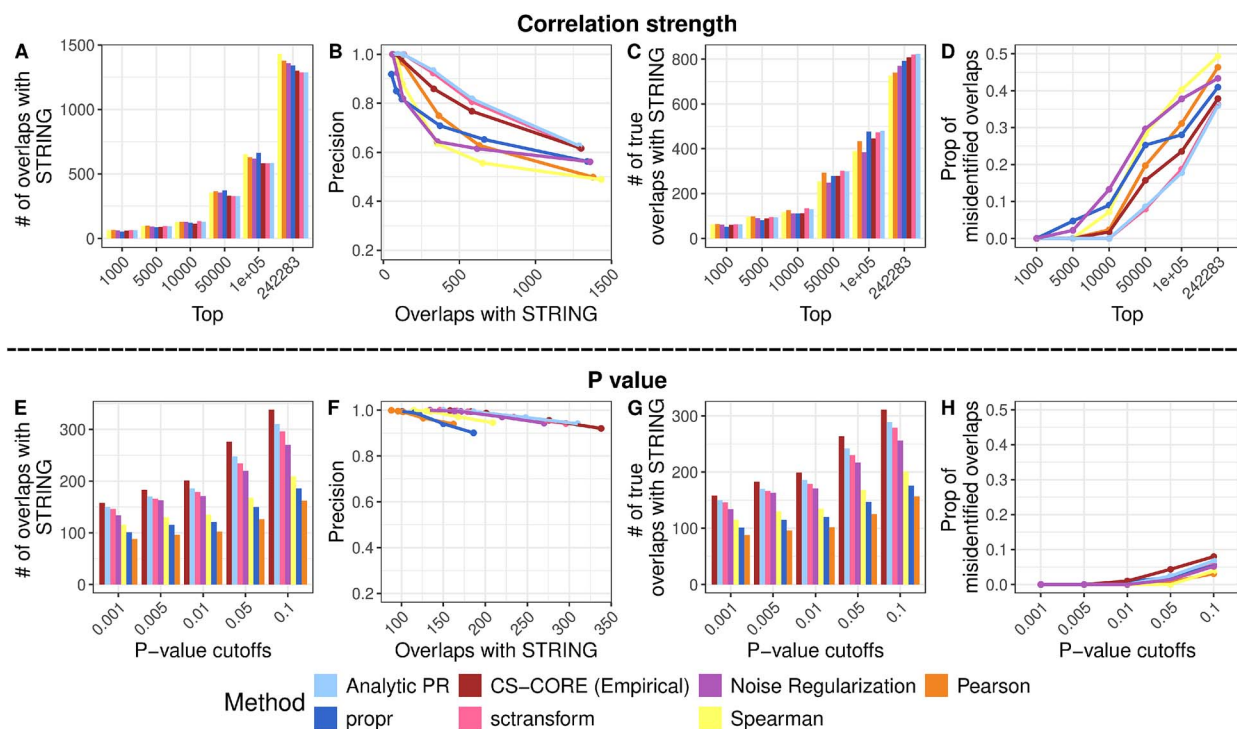


Figure 5 Performance in terms of overlaps with STRING. A and E. The number of overlaps with STRING at different cutoffs. B and F. Relationship between precision and the number of overlaps with STRING. C and G. The number of true overlaps with STRING at different cutoffs. D and H. The proportion of misidentified overlaps when gene pairs were selected by their estimated correlations (D) or by the empirical p -values (H).

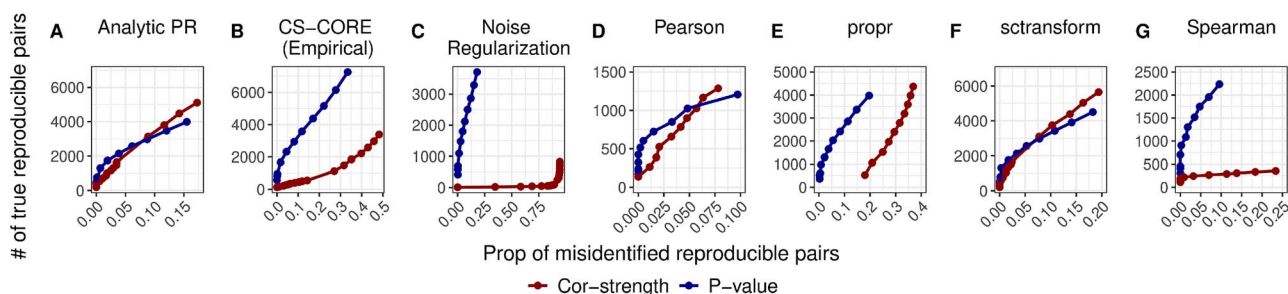


Figure 6 Comparison of correlated gene pair selection approaches in terms of reproducibility. The performance of selecting gene pairs by correlation strength (dark red) versus p -value (dark blue) was evaluated for seven different co-expression estimation methods. Each panel plots the number of true reproducible pairs (y-axis) against the proportion of misidentified reproducible pairs (x-axis) at varying selection thresholds.

quantifies the significance of the observed overlap, corrected for both the total number of identified pairs and expression bias.

Our adjustment strategy is designed to correct for two distinct sources of bias. Therefore, we demonstrate its importance in a two-step comparison. First, we compared the observed overlap counts of different methods with their calculated non-stratified Z-scores, which were generated without adjusting for expression level bias. As shown in Figs 8C and 8F, a naive comparison based on observed overlap counts can be misleading. For example, while CS-CORE with empirical p -values identifies the highest number of overlaps with known biological networks, its Z-score is not among the *top*-ranked. This occurs because a larger set of identified gene pairs inherently increases the chance of random overlaps, confirming that an adjustment for network size is crucial for a fair comparison.

We next demonstrate the necessity of our stratified approach for correcting expression-level bias by comparing its results to those from a non-stratified Z-score. When there is no apparent expression bias in a dataset, both the stratified and non-stratified approaches produce

consistent method significance and rankings (Figs S17A and S17B). However, in the presence of strong expression bias, the stratification step becomes critical for an accurate evaluation. Applying our expression level adjustment has a significant impact on the evaluation outcomes. As shown in Fig. 8F, the Z-scores for some methods become statistically insignificant, and the overall ranking of the methods changes. This demonstrates that the stratified approach is essential for disentangling a method's true performance from the confounding effects of expression bias in validation databases. Furthermore, we found that this adjustment is robust to the choice of interval size (Fig. S18).

The flexibility of our stratified framework allows it to be applied to other evaluation metrics, such as reproducibility. To demonstrate this, we compared the observed number of reproducible pairs between the ROSMAP and PNAS datasets with the Z-scores. The results again show that the raw counts can be misleading. While the Spearman method identifies a high number of reproducible pairs, its Z-score is comparable to that of other methods (Fig. S17C). Furthermore, some

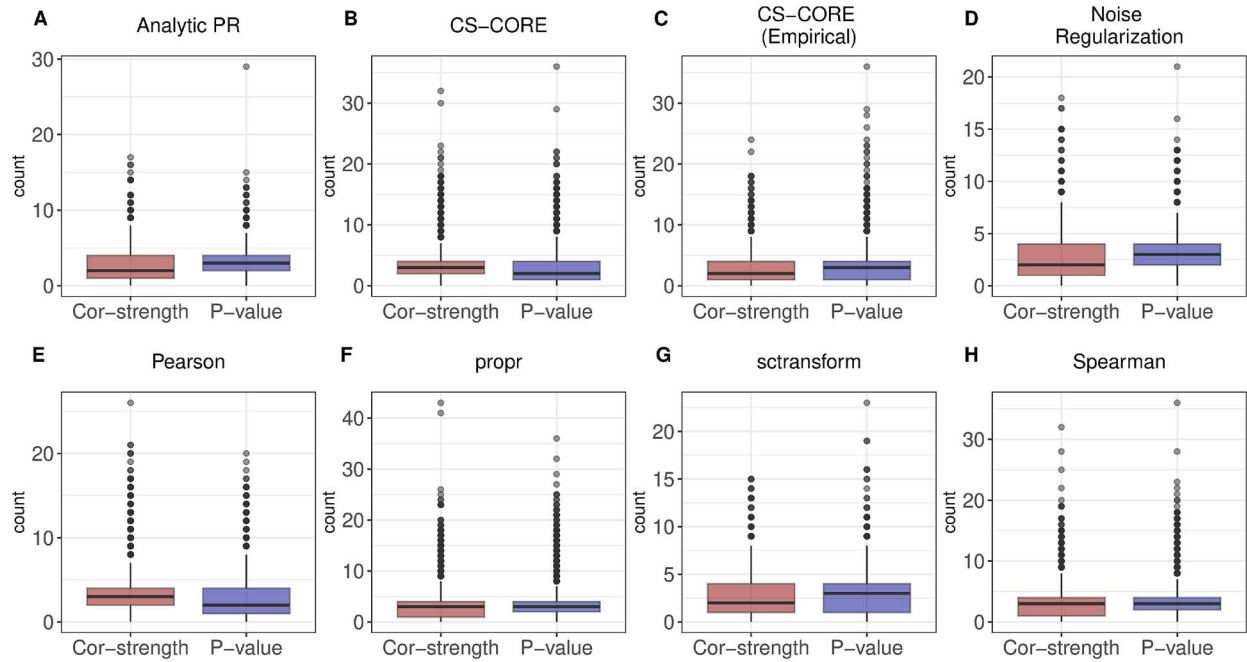


Figure 7 Comparison of shared GO terms for gene pairs identified exclusively by p -value or correlation-strength approaches. For each estimation method, we identified two sets of gene pairs in excitatory neurons, those found only by the p -value-based approach and those found only by the correlation-strength-based approach. To ensure a fair comparison, the threshold for the correlation-strength-based approach was set to yield the same number of gene pairs as the p -value-based approach at a 0.05 significance level (BH adjusted). The boxplots display the distribution of the number of shared GO terms for the gene pairs in each exclusive set.

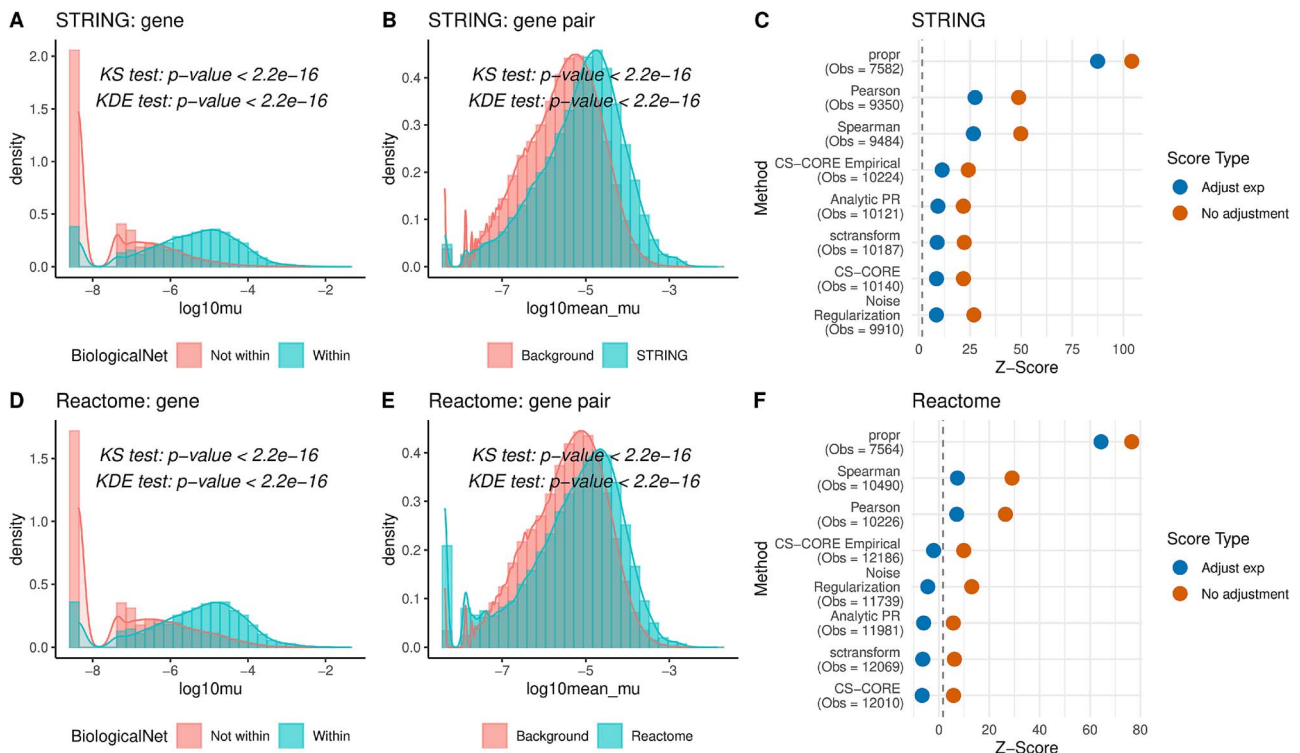


Figure 8 Expression bias in biological databases influences method performance rankings. We evaluated the expression bias in STRING (A–C) and Reactome (D–F). All analyses are based on monocytes from the lung dataset. A and D. Density plots comparing the mean expression of genes present in each network versus those not in the network. KS test and KDE test were conducted to assess the similarity between the two distributions, which demonstrate a significant gene-level bias. B and E. Density plots comparing the mean expression of gene pairs present in each network versus a background set, demonstrating a significant pair-level bias. C and F. Dot plots showing the Z-scores of various co-expression estimation methods with (Adjust exp) and without (No adjustment) the stratified expression adjustment. The number of observed overlaps with the biological database is shown in parentheses (Obs) for each method. The dashed line represents the significance threshold (one-sided $p < 0.05$), which corresponds to a Z-score of 1.645.

co-expression estimation methods are inherently biased towards finding correlations among highly-expressed genes (Fig. S19). Our stratified adjustment corrects for this by evaluating performance within different expression bins, providing a much fairer test. As seen with the Spearman that has a more obvious expression bias (Fig. S19), its Z-score drops more after our adjustment is applied (Fig. S17C).

Discussion

In this study, we first established that, for cell-type-specific analysis of scRNA-seq data, selecting correlated gene pairs based on p -values is a more robust and biologically meaningful approach than relying on correlation strength. This is because p -values, unlike correlation strengths, provide a statistical measure of confidence that accounts for both the effect size and the variability of the data. Our initial results confirm that a p -value-based selection consistently yields a higher proportion of true positives for CS-CORE. However, the practical application of this is limited, as many co-expression estimation methods either do not provide p -values or provide ones that are not well-controlled. To overcome this, we proposed scSimu, an accurate and efficient simulation method that can generate reliable empirical p -values for any co-expression estimation methods. This enabled us to demonstrate the superiority of the p -value-based approach across all estimation methods, and provide a universal and accessible framework to move beyond simple correlation strength and identify more robust correlated gene pairs.

Building on this p -value-based foundation, we next addressed the challenges in the fair method comparison. A simple comparison of overlaps is often misleading, as a method that identifies a larger number of gene pairs will have a higher chance of random overlaps. Furthermore, we illustrated the potential expression bias in the validation databases. To correct for both of these confounding factors, we proposed a stratified hypergeometric model to ensure that the methods are evaluated by their true ability to detect co-expressions, rather than artifacts.

However, some limitations should be noted. First, the validation of scSimu was primarily performed on the droplet-based scRNA-seq data. Although we demonstrated its applicability to plate-based data like Smart-seq (Fig. S21), a more extensive benchmark across a wider range of simulation methods, protocols and datasets would further evaluate its generalizability. Second, our empirical p -value approach, while robust, is computationally expensive (More details can be found in Fig. S22). This may make it less practical for co-expression methods that already have extensive running times, for example, baredSC [27] and locCSN [28]. Third, our framework is not compatible with methods that rely on quantile normalization, such as SpQN [7], as this transformation alters the data distribution in a way that makes the simulated null data incomparable to the observed data. Finally, while our stratified model mitigates the impact of expression bias in validation databases, we still recommend considering diverse evaluation metrics as it is still unknown whether this bias is an artifact or represents real biology.

Ultimately, our study provides a rigorous workflow for co-expression analysis (Fig. S23), moving from the initial principle of p -value superiority to a comprehensive framework for robust gene pair selection and unbiased method benchmarking. By providing both a practical guide and a new standard for evaluation, we hope

to improve the reliability of co-expression analysis in the single-cell genomics field.

Methods

Datasets and preprocessing

Our study utilized three droplet-based datasets, including two single-nucleus RNA-sequencing (snRNA-seq) datasets of the human prefrontal cortex from the ROSMAP cohort [29] and a study by Lau *et al.* [30] (PNAS), and one scRNA-seq dataset from lung tissue [26]. To ensure consistency, we restricted our analysis to control subjects from each study. Detailed preprocessing steps are provided in [Supplementary Note S1](#).

Biological networks. We obtained the human protein-protein interaction network from STRINGv12 [13], filtering for interactions with a combined score greater than 500. The other biological network we considered was a pathway database, Reactome [14]. By combining each pair within the same pathway, we constructed an interaction network [1]. We also utilized the CollecTRI database [31] to generate a cell-type-specific TF-target reference by subsetting for oligodendrocyte marker genes identified from ROSMAP single-cell data, with additional details provided in [Supplementary Note S1](#).

GO terms. The genes' GO terms were obtained from the R package biomaRt version 2.58.0.

scSimu

For cell i and gene j , scSimu simulates the observed count, x_{ij} , from the following distributions,

$$\begin{aligned} z_{ij} &\sim F_j = \text{Gamma}(\alpha_j, \mu_j), \\ x_{ij}|z_{ij} &\sim \text{Poisson}(s_i z_{ij}), \end{aligned} \quad (1)$$

where z_{ij} is the corresponding relative underlying expression, α_j is the dispersion parameter, μ_j is the mean expression, and s_i is the sequencing depth. We model the dependence structure between genes at the relative underlying expression level using a copula, which can introduce the dependence among genes while preserving the marginal distributions. Based on Sklar's theorem, the joint cumulative distribution function (CDF) of the relative underlying expression for p genes can be written as $F(z_{i1}, \dots, z_{ip}) = C(F_1(z_{i1}), \dots, F_p(z_{ip}))$, where $C(\cdot)$ is a copula. By employing a Gaussian copula with a target gene correlation matrix R , this relationship is defined specifically as $F(z_{i1}, \dots, z_{ip}) = \Phi_R(\Phi^{-1}(F_1(z_{i1})), \dots, \Phi^{-1}(F_p(z_{ip})))$, where Φ_R is the CDF of multivariate normal distribution with mean vector zero and covariance matrix equal to the correlation matrix R , and Φ^{-1} is the inverse CDF of standard normal distribution. Therefore, we can use an inverse procedure to generate a multivariate distribution with the specified correlation structure and marginal distributions.

Compared with the simulation method used in our previous publication [11], here we updated the Gaussian copula process to allow for the case that the correlation matrix is not positive semi-definite given that some gene co-expression estimation methods sacrifice positive semi-definite to accommodate the noise and sparsity in scRNA-seq (e.g. CS-CORE, SpQN). Additionally, after applying a threshold based on significance or correlation strength, a positive semi-definite correlation matrix may no longer remain positive semi-definite. Our updated Gaussian copula process addresses this by employing eigendecomposition to replace any negative eigenvalues with zero. This procedure

provides a simulation with a correlation structure that is more similar to the true one compared to our previous method of iteratively adding small values to the diagonal (Fig. S20). In addition, unlike methods such as SymSim [32], ZINB-WaVE, powsimR [33], and muscat [34], scSimu preserves all original genes and cells in the output, which is critical for accurate downstream comparisons.

The detailed simulation procedures and the zero-inflated model are described in [Supplementary Note S2](#).

Benchmark simulation methods

To evaluate the performance of different simulation methods, we compared the simulated data with real data from eight data summaries:

- **Mean(logCPM)**: mean expression of genes estimated via log-normalized data.
- **Var(logCPM)**: variance of genes estimated using log-normalized data.
- **Coefficient of variation**: $\sqrt{\text{Var}(\log\text{CPM})}/\text{Mean}(\log\text{CPM})$.
- **Zero fraction (cell)**: fraction of cells with zero counts.
- **Zero fraction (gene)**: fraction of genes with zero counts.
- **Library size**: total gene counts per cell.
- **Cell correlation**: Spearman correlation between cells based on log-normalized data. The number of cells used to estimate cell correlation was capped by 500.
- **Gene correlation**: correlations estimated by CS-CORE for the top 1,000 highly expressed genes determined by Mean(logCPM).

To quantify the similarity between simulated and real data distributions, we employed both Kolmogorov-Smirnov (KS) and Kernel Density Estimation (KDE) tests. These tests offer complementary insights, with the KS test measuring the largest discrepancy between distributions and the KDE test assessing overall divergence. To ensure comparability across different datasets and metrics, KS and KDE statistics were min-max normalized separately by dataset and metric, then averaged to combine them. The overall performance of each simulation method was evaluated based on the mean score across eight data metrics. Comprehensive details and the benchmark for Smart-Seq data simulation are presented in [Supplementary Note S3](#).

Empirical p -value estimation

The empirical p -values were estimated in three steps. (1) Generate k independent datasets corresponding to the real data using scSimu and estimate the correlations. For each gene pair, gather k estimates from these independent datasets alongside one from the real data. (2) Estimate the mean μ_0 and variance σ_0^2 of these k estimates for each gene pair. (3) The empirical p -value is calculated using a normal distribution with the estimated mean and variance, where the p -value = $P(|X| > |Estimate_{obv}|)$ and $X \sim N(\mu_0, \sigma_0^2)$. See [Supplementary Note S4](#) for the analysis of type I error control.

Evaluation of estimated gene co-expression networks

To compare the performance of p -value-based and correlation-strength-based approaches, we generated synthetic datasets with known ground truth based on oligodendrocytes from ROSMAP and PNAS control subjects. We created a sparse ground truth by setting weak correlations to zero in the real data and simulating new counts using scSimu. We then estimated co-expression networks for the top 1,000 highly expressed shared genes using seven different methods. More details can be found in [Supplementary Note S5](#).

Reproducibility. We used the simulated PNAS data as independent validation data. The reproducibility was evaluated based on the number of reproducible pairs, which are gene pairs identified as correlated in both datasets.

Precision. Precision is the ratio of true positives to inferred positives. For reproducibility, the precision values in the main text were based on the ROSMAP data, while those from the PNAS data were displayed in [Figs S1 and S6](#).

Proportion of misidentification. Let A and B be the sets of identified correlated gene pairs from two datasets, with A_T and B_T representing their respective true positives. The proportion of misidentified reproducible pairs is calculated as $(|A \cap B| - |A_T \cap B_T|)/|A \cap B|$. Similarly, for a biological network C , the proportion of misidentified overlaps is $(|A \cap C| - |A_T \cap C|)/|A \cap C|$.

Expression bias in biological networks

To demonstrate the expression bias in the STRING and Reactome, we estimated the mean expression level for each gene using monocytes from the lung dataset and a Gamma-Poisson GLM. We first assessed the gene-level bias by comparing the expression distributions of genes in these biological networks versus those that are not. We then focused on genes within the networks to evaluate pair-level bias. We further investigated the bias by repeating the pair-level analysis on subsets of genes, including the top 100, 1,000, and 5,000 most highly expressed or highly variable genes. For each of these analyses, the biological network was correspondingly filtered to include only interactions among the selected top k genes. The background was generated based on this sub-network. To confirm that these findings were not dataset-specific, we also studied the observed bias using mean expression levels from oligodendrocytes in the ROSMAP.

A stratified hypergeometric model for bias adjustment.

First, all gene pairs are categorized into k bins based on their mean expression level. The number of random overlaps is then modeled by a separate hypergeometric distribution within each bin i , $\text{Hypergeometric}(n_i, b_i, c_i)$. The total expected number of random overlaps and its variance are calculated by summing the expectations and variances from all k bins, $E(\text{Random overlaps}) = \sum_{i=1}^k \frac{c_i b_i}{n_i}$ and $\text{Var}(\text{Random overlaps}) = \sum_{i=1}^k \frac{c_i b_i}{n_i} \frac{n_i - b_i}{n_i} \frac{n_i - c_i}{n_i - 1}$. This yields a stratified Z-score corrected for both network size and expression bias. We validated the stratified framework using the lung dataset by assessing the overlap of identified networks with the STRING and Reactome databases. Details regarding the baseline non-stratified model and the application to real data are provided in [Supplementary Note S6](#).

Key Points

- Our results demonstrate that selecting correlated gene pairs based on p -values is more robust and biologically meaningful than relying on the commonly used correlation strength.
- We reveal expression bias in standard biological validation databases and propose a stratified hypergeometric model to correct for this confounding factor during benchmarking.

- This study improves the reliability of co-expression analysis by establishing a comprehensive workflow for unbiased method evaluation and reliable gene pair selection.

Acknowledgements

We thank Dr Jia Zhao for her insightful feedback and helpful discussions. We also would like to thank the Religious Orders Study and Memory and Aging Project (ROSMAP) Study, Dr Shun-Fat Lau, Dr Nancy Y. Ip, Taylor S. Adams, Dr Naftali Kaminski, Dr Ding Jiarui, and Dr Joshua Z. Levin for sharing the datasets.

Authors' contributions

X.S. conducted the study. Y.L. participated the discussion of findings and provided feedback on the manuscript draft. H.Z. supervised the project. All authors read and approved the final manuscript.

Supplementary data

Supplementary data is available at *Briefings in Bioinformatics* online.

Conflicts of interest

None declared.

Funding

This work was supported in part by NIH grants R01 GM134005, U01 HG013840, and U24 HG012108.

Availability of data and materials

The ROSMAP dataset (syn52293417) has controlled access under the regulation of human privacy. The PNAS dataset used in our analysis is publicly available ([GSE157827](https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GSE157827)). The lung data we used to evaluate the expression bias in biological networks can be accessed through [GSE136831](https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GSE136831). The Smart-Seq data can be accessed through the Single Cell Portal with accession numbers [SCP424](https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=SCP424). The relevant code of this study can be found in https://github.com/xs222/coexp_validation [35]. The R package scSimu is available at <https://github.com/xs222/scSimu> [36].

References

1. Skinnider MA, Squair JW, Foster LJ. Evaluating measures of association for single-cell transcriptomics. *Nat Methods* 2019;**16**:381–6. <https://doi.org/10.1038/s41592-019-0372-4>
2. Li WV, Li Y. Sclink: Inferring sparse gene co-expression networks from single-cell expression data. *Genom Proteom Bioinform* 2021;**19**:475–92. <https://doi.org/10.1016/j.gpb.2020.11.006>
3. Iacono G, Massoni-Badosa R, Heyn H. Single-cell transcriptomics unveils gene regulatory network plasticity. *Genome Biol* 2019;**20**:1–20. <https://doi.org/10.1186/s13059-019-1713-4>
4. Pratapa A, Jalihal AP, Law JN *et al*. Benchmarking algorithms for gene regulatory network inference from single-cell transcriptomic data. *Nat Methods* 2020;**17**:147–54. <https://doi.org/10.1038/s41592-019-0690-6>
5. Crow M, Paul A, Ballouz S *et al*. Exploiting single-cell expression to characterize co-expression replicability. *Genome Biol* 2016;**17**:1–19. <https://doi.org/10.1186/s13059-016-0964-6>
6. Zhang R, Atwal GS, Lim WK. Noise regularization removes correlation artifacts in single-cell rna-seq data preprocessing. *Patterns* 2021;**2**:100211. <https://doi.org/10.1016/j.patter.2021.100211>
7. Wang Y, Hicks SC, Hansen KD. Addressing the mean-correlation relationship in co-expression analysis. *PLoS Comput Biol* 2022;**18**:1009954. <https://doi.org/10.1371/journal.pcbi.1009954>
8. Lu S, Keleş S. Debiased personalized gene coexpression networks for population-scale scRNA-seq data. *Genome Res* 2023;**33**:932–47. <https://doi.org/10.1101/gr.277363.122>
9. Quinn TP, Richardson MF, Lovell D *et al*. Propr: an R-package for identifying proportionally abundant features using compositional data analysis. *Sci Rep* 2017;**7**:1–9. <https://doi.org/10.1038/s41598-017-16520-0>
10. Laue J, Berens P, Kobak D. Analytic Pearson residuals for normalization of single-cell rna-seq umi data. *Genome Biol* 2021;**22**:1–20. <https://doi.org/10.1186/s13059-021-02451-7>
11. Su C, Xu Z, Shan X *et al*. Cell-type-specific co-expression inference from single cell rna-sequencing data. *Nat Commun* 2023;**14**:4846. <https://doi.org/10.1038/s41467-023-40503-7>
12. Li S, Schmid KT, Vries DH *et al*. Identification of genetic variants that impact gene co-expression relationships using large-scale single-cell data. *Genome Biol* 2023;**24**:1–37. <https://doi.org/10.1186/s13059-023-02897-x>
13. STRING: Functional Protein Association Networks. Accessed: 2024-08-02. <https://string-db.org/>.
14. Reactome Pathway Database. Accessed: 2024-08-02. <https://reactome.org/>.
15. Chen S, Mar JC. Evaluating methods of inferring gene regulatory networks highlights their lack of performance for single cell gene expression data. *BMC Bioinform* 2018;**19**:1–21. <https://doi.org/10.1186/s12859-018-2217-z>
16. Wang L. Single-cell normalization and association testing unifying crispr screen and gene co-expression analyses with normalizr. *Nat Commun* 2021;**12**:6395. <https://doi.org/10.1038/s41467-021-26682-1>
17. Hafemeister C, Satija R. Normalization and variance stabilization of single-cell rna-seq data using regularized negative binomial regression. *Genome Biol* 2019;**20**:296. <https://doi.org/10.1186/s13059-019-1874-1>
18. Crowell HL, Morillo Leonardo SX, Sonesson C *et al*. The shaky foundations of simulating single-cell rna sequencing data. *Genome Biol* 2023;**24**:1–19. <https://doi.org/10.1186/s13059-023-02904-1>
19. Risso D, Perraudeau F, Gribkova S *et al*. A general and flexible method for signal extraction from single-cell rna-seq data. *Nat Commun* 2018;**9**:284. <https://doi.org/10.1038/s41467-017-02554-5>
20. Sun T, Song D, Li WV *et al*. scdesign2: a transparent simulator that generates high-fidelity single-cell gene expression count data with gene correlations captured. *Genome Biol* 2021;**22**:163. <https://doi.org/10.1186/s13059-021-02367-2>
21. Baruzzo G, Patuzzi I, Di Camillo B. Sparsim single cell: a count data simulator for scRNA-seq data. *Bioinformatics* 2020;**36**:1468–75. <https://doi.org/10.1093/bioinformatics/btz752>
22. Cao Y, Yang P, Yang JYH. A benchmark study of simulation methods for single-cell rna sequencing data. *Nat Commun* 2021;**12**:6911. <https://doi.org/10.1038/s41467-021-27130-w>
23. Song D, Wang Q, Yan G *et al*. scdesign3 generates realistic in silico data for multimodal single-cell and spatial omics. *Nat Biotechnol* 2023;**42**:247–52. <https://doi.org/10.1038/s41587-023-01772-1>

24. Jiang R, Sun T, Song D *et al.* Statistics or biology: the zero-inflation controversy about scRNA-seq data. *Genome Biol* 2022;**23**:1–24. <https://doi.org/10.1186/s13059-022-02601-5>
25. Yuan L, Xu Z, Meng B *et al.* Scamzi: attention-based deep autoencoder with zero-inflated layer for clustering scRNA-seq data. *BMC Genom* 2025;**26**:350. <https://doi.org/10.1186/s12864-025-11511-2>
26. Adams TS, Schupp JC, Poli S *et al.* Single-cell RNA-seq reveals ectopic and aberrant lung-resident cell populations in idiopathic pulmonary fibrosis. *Sci Adv* 2020;**6**:1983. <https://doi.org/10.1126/sciadv.aba1983>
27. Lopez-Delisle L, Delisle J-B. Baredsc: Bayesian approach to retrieve expression distribution of single-cell data. *BMC Bioinform* 2022;**23**:36. <https://doi.org/10.1186/s12859-021-04507-8>
28. Wang X, Choi D, Roeder K. Constructing local cell-specific networks from single-cell data. *Proc Natl Acad Sci* 2021;**118**:2113178118. <https://doi.org/10.1073/pnas.2113178118>
29. Mathys H, Peng Z, Boix CA *et al.* Single-cell atlas reveals correlates of high cognitive function, dementia, and resilience to Alzheimer's disease pathology. *Cell* 2023;**186**:4365–4385.e27. <https://doi.org/10.1016/j.cell.2023.08.039>
30. Lau S-F, Cao H, Fu AK *et al.* Single-nucleus transcriptome analysis reveals dysregulation of angiogenic endothelial cells and neuroprotective glia in Alzheimer's disease. *Proc Natl Acad Sci* 2020;**117**:25800–9. <https://doi.org/10.1073/pnas.20088&break;762117>
31. Müller-Dott S, Tsirvouli E, Vazquez M *et al.* Expanding the coverage of regulons from high-confidence prior knowledge for accurate estimation of transcription factor activities. *Nucleic Acids Res* 2023;**51**:10934–49. <https://doi.org/10.1093/nar/gkad841>
32. Zhang X, Xu C, Yosef N. Simulating multiple faceted variability in single cell RNA sequencing. *Nat Commun* 2019;**10**:2611. <https://doi.org/10.1038/s41467-019-10500-w>
33. Vieth B, Ziegenhain C, Parekh S *et al.* Powsimr: power analysis for bulk and single cell RNA-seq experiments. *Bioinformatics* 2017;**33**:3486–8. <https://doi.org/10.1093/bioinformatics/btx435>
34. Crowell HL, Sonesson C, Germain P-L *et al.* Muscat detects subpopulation-specific state transitions from multi-sample multi-condition single-cell transcriptomics data. *Nat Commun* 2020;**11**:6077. <https://doi.org/10.1038/s41467-020-19894-4>
35. Shan X, Zhao H. *Code for Coexpression Validation*. Accessed: 2025-10-01 (2025). https://github.com/xs222/coexp_validation.
36. Shan X, Zhao H. *scSimu: A Simulation Tool for Single-Cell RNA Sequencing Data*. Accessed: 2025-10-01 (2025). <https://github.com/xs222/scSimu>.