

Refining the resolution of the yeast genotype-phenotype map using single-cell RNA-sequencing

Reviewed Preprint

v2 • September 20, 2024

Revised by authors

Reviewed Preprint

v1 • March 20, 2024

Arnaud N'Guessan, Wen Yuan Tong, Hamed Heydari, Alex N Nguyen Ba 

Department of Cell and Systems Biology, University of Toronto, Ramsay Wright Laboratories, 25 Harbord St, M5S3G5, Toronto, Ontario, Canada • Department of Biology, University of Toronto at Mississauga, 3359 Mississauga Rd, L5L 1C5, Mississauga, Ontario, Canada • Department of Molecular Genetics, University of Toronto, Medical Science Building, Room 4386, 1 King's College Cir, M5S1A8, Toronto, Ontario, Canada • Donnelly Centre for Cellular and Biomolecular Research, University of Toronto, 160 College Street, Room 230, M5S3E1, Toronto, Ontario, Canada

 https://en.wikipedia.org/wiki/Open_access

 Copyright information

Abstract

Genotype-phenotype mapping (GPM) or the association of trait variation to genetic variation has been a long-lasting problem in biology. The existing approaches to this problem allowed researchers to partially understand within- and between-species variation as well as the emergence or evolution of phenotypes. However, traditional GPM methods typically ignore the transcriptome or have low statistical power due to challenges related to dataset scale. Thus, it is not clear to what extent selection modulates transcriptomes and whether cis- or trans-regulatory elements are more important. To overcome these challenges, we leveraged the cost efficiency and scalability of single-cell RNA sequencing (scRNA-seq) by collecting data from 18,233 yeast cells from 4,489 F2 segregants derived from an F1 cross between the laboratory strain BY4741 and the vineyard strain RM11-1a. More precisely, we performed eQTL mapping with the scRNA-seq data to identify single-cell eQTL (sc-eQTL) and transcriptome variation patterns associated with fitness variation inferred from the segregant bulk fitness assay. Due to the larger scale of our dataset and its multidimensionality, we could recapitulate results from decades of work in GPM from yeast bulk assays while revealing new associations between phenotypic and transcriptomic variations at a broad scale. We evaluated the strength of the association between phenotype variation and expression variation, revealed new hotspots of gene expression regulation associated to trait variation, revealed new gene function with high expression heritability and highlighted the larger aggregate effect of trans-regulation compared to cis-regulation. Altogether these results suggest that integrating large-scale scRNA-seq data into GPM improves our understanding of trait variation in the context of transcriptomic regulation.

eLife assessment

This **useful** study describes expression profiling by scRNA-seq of thousands of cells of recombinant yeast genotypes from a system that models natural genetic variation. The rigorous new method presented here shows promise for improving the efficiency of genotype-to-phenotype mapping in yeast, providing **convincing** evidence for its efficacy. This revised manuscript focuses on overcoming technical challenges with this approach and identifies several new biological insights that build upon the field of genotype-to-phenotype mapping, a central question of interest to geneticists and evolutionary biologists.

<https://doi.org/10.7554/eLife.93906.2.sa3>

Introduction

The process by which DNA encodes proteins via transcription and translation has been studied for decades to make sense of organisms' phenotypes. However, being able to explain organisms' phenotypes from their genetic material, i.e. genotype-phenotype mapping (GPM), has been a long-lasting problem with important applications (1,2). Indeed, making sense of genetic variation at the phenotypic level enables the understanding of trait variation between and within species as well as the emergence and evolution of phenotypes (3). For instance, reverse genetics approaches, e.g. gene knockout or transgenic technologies, and forward genetics approaches like GWAS and QTL mapping helped in determining the function of multiple genes and the effects of mutations on growth in different environments (4). However, reverse genetics approaches typically fail to account for natural variation and forward genetics approaches like QTL mapping typically focus on genetic and phenotypic variation so they cannot highlight selection on the transcriptome.

An essential characteristic of this problem is the multi-layered organization of the GPM. Indeed, GPM is not strictly restricted to the direct association between genotypes and phenotypes. This association is better resolved and complemented by understanding the intermediary transcriptome layer, e.g. cell mechanisms at the transcriptomic level are involved in diseases and pathogenicity (2,5,6–8). However, it is not clear to what extent transcriptomic changes relate to phenotypic changes or selection. Pioneering work from Mary-Claire King and Allan Charles Wilson set the tone for investigating this question by proposing that variations in morphological and behavioral traits arise more often through gene expression regulation than evolution at the protein-coding level (9). François Jacob then postulated an essay that stemmed from this theory in which he highlights how evolution acts as a tinkerer that works from already available material, i.e. through regulation of gene expression, to create new adaptations (10). This constituted the core of the evolutionary developmental biology which matured into the still-debated claim that new adaptations mainly emerge through cis-regulation of gene expression, i.e. through noncoding DNA regulating a neighbor gene contrarily to trans-regulators acting on distant genes (11–14). This debate has been reinforced by the technical difficulties and complexity of assessing the evolution and outcome of mutations in non-coding regions (11,12). For example, although human genome-wide association studies (GWAS) loci are mostly identified in noncoding regions and have well-characterized cis-regulatory effects, their trans-regulatory effects are usually challenging to detect (15–17). Indeed, small effect sizes, multiple testing corrections and high false-positive rates for trans-regulation impose statistical and computational challenges for its detection (18,19). Advances in sequencing technologies helped to partially solve these issues, particularly in the context of transcriptome analyses of the model organism

Saccharomyces cerevisiae. For instance, Brem et al (2002) used microarray technology to relate the gene expression profiles of 40 yeast segregants from a lab (BY) and natural vineyard strain (RM) to their genetic markers (20). They found that cis-acting modulation is the main mechanism for regulating gene expression. Nearly two decades later, by greatly increasing statistical power, Albert and collaborators (2018) found that most of the expression variation arises through trans-regulation using non-multiplexed RNA-seq to analyze 5,720 genes in 1,012 yeast segregants generated by a crossing between RM and BY (21). The analysis method they used, i.e. expression quantitative trait loci (eQTL) mapping, consists in correlating allele frequencies to gene expression levels to find the loci modulating expression.

Although eQTL mapping is a traditional GPM analysis that accounts for the transcriptomic layer, it is typically realized through non-multiplexed RNA-seq which tends to have low statistical power due to challenges with experimental scale and confounding factors (22,23). Thus, eQTL mapping traditionally cannot identify significant low-effect regulatory mutations that are important for understanding the genetic bases of complex traits and diseases (24,25). Furthermore, most eQTL studies only assess the average transcriptomic profile of bulk populations without being able to capture the profile of rare cell lineages within a population. This is a critical limitation in heterogeneous populations such as cancer or microbial populations where rare lineages can drive relapse or drug resistance (26).

Here, we sought to circumvent the challenges of non-multiplexed bulk RNA-seq imposed by the scale and population heterogeneity by performing eQTL mapping through single-cell RNA sequencing (scRNA-seq) of a pool of ~4,500 well-characterized F2 segregants derived from a yeast cross (21,27,28). In the same way that combinatorial indexing/barcoding and multiplexing enable the collection of large-scale fitness and genotype data (29), we hypothesized that scRNA-seq could help us collect both genotype and expression data on a large pool of segregants. We employ several strategies to overcome previous obstacles of eQTL mapping studies: i) we pool cells from thousands of F2 segregants during the growth step and perform a single scRNA-seq run on the culture to account for environmental effects. The expression profiles of these lineages cannot be captured cost-effectively with bulk RNA-seq as pooling segregants in a single bulk assay would not capture individual strains' expression profiles. Moreover, ii) from the exome sequencing data of single cells, we take advantage of the reference panel to validate that we accurately infer the genotype of each cell from extremely low number of reads mapping to polymorphic sites per cell (effectively ~0.2x coverage). Finally, iii) we leverage the presence of thousands of high-coverage cells to train an unsupervised learning model to correct (denoise) and infer an imputed expression profile of poorly covered cells.

Using this approach, we integrated the resulting transcriptomic data from growth in rich media with a pre-existing yeast GPM to reveal new GPM features at a broad scale. For instance, we estimated the heritability of the transcriptome and the extent to which the transcriptome is associated with fitness. This allowed us to reveal that a negligible portion of expression variation is related to the studied phenotypic variation independently of genotypic variation. This implies that almost all the phenotypic variation related to expression is also modulated by mutations. This differentiates our work from previous bulk RNA-seq studies where expression is systematically controlled for differences in lineage growth or phenotypes due to an assumption that a considerable amount of phenotypic variation is related to expression independently of genomic mutations (21). We also show that the increased scale of our scRNA-seq dataset enables single-cell eQTL mapping (sc-eQTL) and identifies more hotspots of expression regulation containing previously identified QTL. We also exploit the identified sc-eQTL to analyze the patterns of cis- and trans-regulation in the GPM.

Our single-cell RNA-seq approach is consistent with yeast GPM results from non-multiplexed assays

We initially aimed to show that performing scRNA-seq at a large scale can generate data that are consistent with non-multiplexed DNA and RNA sequencing. For this purpose, we analyzed a dataset of thousands of yeast lineages generated by Nguyen Ba and collaborators (2022) (29). To understand the yeast GPM, they collected fitness and genotype data from ~100,000 F2 segregants derived from an F1 cross between a laboratory strain of yeast (BY) and a natural vineyard strain (RM) (Figure 1A).

Using this approach named barcoded bulk QTL mapping or BB-QTL mapping, they revealed the complex polygenic and pleiotropic nature of phenotypes as well as an unprecedented number of pairwise epistatic interactions. To integrate transcriptomic data into this GPM, we performed scRNA-seq using the 10X Genomics Chromium microfluidics platform (30). Other scRNA-seq methods like Smart-seq2 offer a better breadth of coverage per cell (31). However, due to the higher throughput or number of cells generated by 10x chromium and the different coverage profiles across cells from the same F2 segregant, it is more advantageous to use 10x Chromium to capture transcriptomic variation and reconstruct the reference genomes from single cells (31,32). The throughput of 10x Genomics is typically in the order of 10^4 cells which limits us to a few thousand lineages per library if we want to obtain a reasonable number of cells per lineage. This is why we focus on a single batch of 4,489 F2 segregants. This method allowed us to obtain both genotype and expression profiles from 18,233 cells of the first batch of F2 segregants (Figure 1B). The F2 segregant barcodes are typically not expressed so single cells have their own cell barcode which is included in all amplified reads of the same droplet. Although this short-read scRNA-seq is a high-throughput approach (Table 1), it comes with challenges like gene dropouts and low sequencing depth in some cells caused by technical artifacts of PCR amplification (25,26).

To overcome these challenges, the unique molecular identifiers (UMIs) of the 10X Genomics platform provide a control for PCR amplification biases by quantifying gene expression from unique transcribed molecule counts instead of sequencing read counts (31). However, UMIs still do not correct for dropouts or missing data in the expression profile and this could obscure potential associations between genotype and transcription (see below). We addressed this issue by fitting an imputation model called DISCERN to the expression data (33). DISCERN is a neural network that learns how to reconstruct the expression profile of high-coverage cells after embedding it in a lower dimension. This eliminates the gene dropouts and denoises the expression profile. The fact that the reconstructed expression is highly but not perfectly correlated to the original expression (mean $R^2 = 50.0\%$, median $R^2 = 52.9\%$; R^2 s.d. = 10.1%) and the high variance explained by the first principal component of the imputed expression (99.6%) is consistent with an effective denoising. We take advantage of this denoised transcriptome for variance partitioning (as described later).

In addition, Hidden Markov Models (HMMs) can infer accurate genotype data even at sequencing depths as low as 0.1x (29). Nguyen Ba and collaborators (2022) designed an HMM to infer the segregants genotypes from the observed reads at low depth of DNA sequencing by accounting for sequencing error rate, recombination rate, and index swapping rate (29). As there are only two ancestral lineages, there are only two possible alleles for the strains at each of the 41,594 polymorphic sites. Thus, the genotype of the segregants can be represented by the frequency of only one of the parental alleles, which is RM in the dataset. Applying this model to low-coverage segregants yielded genotypes that are significantly similar to high-coverage replicates (29). We sought to use a similar model to infer genotypes from scRNA-seq data but we anticipated that some of these parameters may differ due to increased error rate of the reverse-transcriptase, increased index swapping due to pooled reaction, etc (Figure S1). In Nguyen Ba et al, those rates

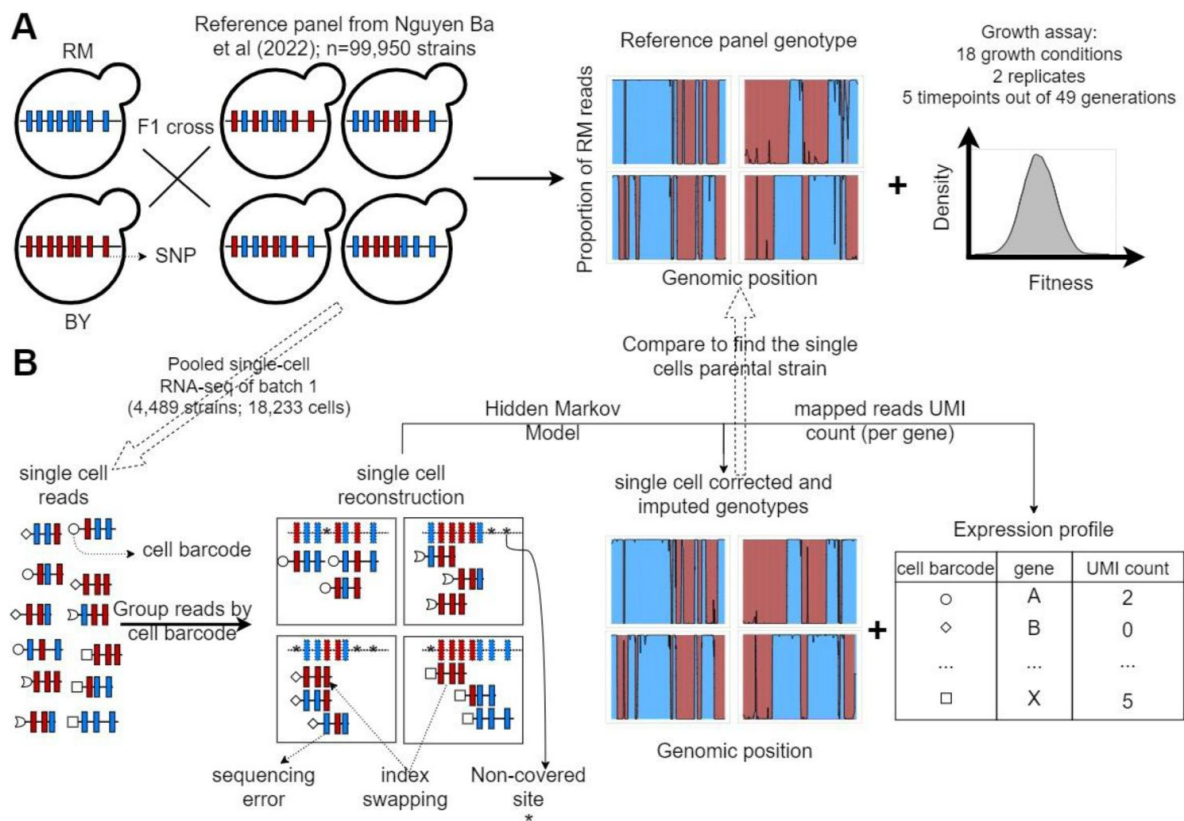


Figure 1

F2 yeast segregants datasets.

A) Reference panel and experiments from the barcoded bulk sequencing previously described in (29). The 99,995 F2 yeast segregants in the reference panel were derived from an F1 cross between a laboratory strain of yeast (BY) and a natural vineyard strain (RM). Thus, they only have 2 possible alleles at each of the 41,594 polymorphic sites. The lineage barcodes enabled fitness estimation from competition assays in 18 environments recapitulating the adaptation to temperature gradients, the ability to process different sources of carbon and the resistance to antifungal compounds. B) Pooled scRNA-seq dataset from a single batch. We performed scRNA-seq of the first batch of F2 segregants (n=4,489) to obtain genotypes that are similar to the reference panel and cell's expression profiles. The F2 segregant barcodes are typically not expressed so single cells have their own cell barcode which is included in all amplified reads of the same droplet. Non-covered sites, sequencing errors and the presence of reads in the wrong library (index swapping) are corrected for using the HMM described in Figure S1.

Metric	Value
Library size (number of reads)	479,781,081
Median number of covered genes per cell	1,311
Median number of UMIs per cell	5,158
Doublet rate	0.10
Average number of cell barcodes per F2 segregant	4.83
Polymorphic sites mean breadth of coverage	0.03

Table 1

scRNA-seq dataset features

. The doublets were defined as scRNA-seq barcodes with high coverage ($\geq$ third quartile) without significant genotype association to any F2 segregants and for which the 2 most similar F2 segregants do not share genotypic similarity ($R^2 < 0.1$). The breadth of coverage is defined as the proportion of BY/RM polymorphic sites covered in a single cell.

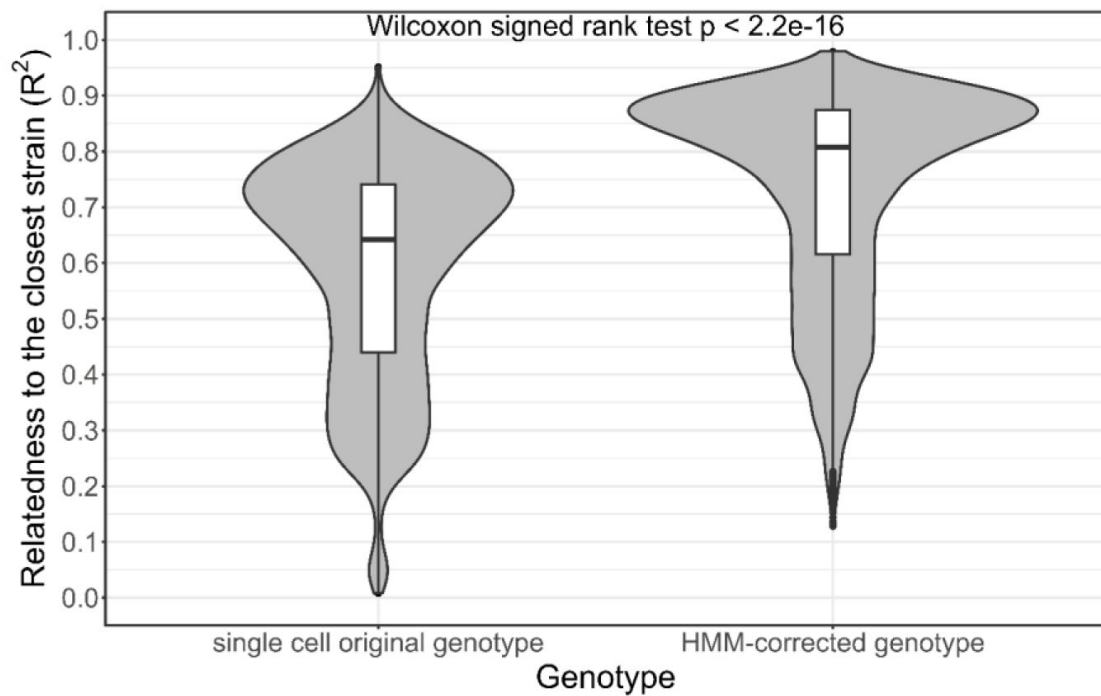
were heuristically determined, but here we estimated these from the read mapping data and found that re-estimated parameters from data increased the proportion of recovered strains in the single cell data from 58.6% to 72.0%.

After adapting the HMM to the scRNA-seq data, we sought to validate that the resulting cell genotypes relate well to their corresponding strain in the reference panel obtained by non-multiplexed DNA sequencing strategies. Ideally, each single-cell barcode (from 10x Genomics Chromium) should be associated with a single cell and a cell should have a clear match with a unique strain in the reference panel. However, several factors can obscure these associations, e.g. a single-cell droplet containing cells from 2 different strains, a low-coverage cell, uncertainty in the allele of the reference genotype, etc. Thus, we designed an approach to assign cells to the correct reference panel strain (see **Methods**). This approach relies on two metrics of similarity between the cells and the strains' genotypes, i.e. the expected distance between them, which should be minimized for the best match, and the relatedness (R^2). The statistical significance of the relatedness between single cells and reference lineages was determined by a permutation test (**Figure S2**). From the read mapping alone, we obtained a mean R^2 of 0.59 ($\sigma = 0.19$ and median = 0.64), which was significantly improved after applying our HMM to correct for misidentified alleles and imputing data in low-coverage sites using recombination probability. Indeed, the single-cell HMM genotypes yield a mean R^2 of 0.73 ($\sigma = 0.18$ and median = 0.81; **Figure 2A**). We found that the distribution of relatedness after HMM was still left-skewed, with many cells statistically significantly assigned to a reference genotype despite having what appeared to be low relatedness. Upon investigation, it was found that these could be explained by genotyping uncertainty either in the single-cell and/or in the reference panel genotype (s) (**Table S1**).

The single-cell original genotype represents the genotype of the cells before the correction with the HMM. The relatedness to the closest lineage in batch 1 has been measured with the adjusted R^2 . To control for genotype uncertainty, only the 13,069 barcodes with a significant lineage assignment (lineage-barcode genotype correlation FDR<0.05) and a reference lineage with lower uncertainty than the single cell HMM are selected, which represents 72.2% of the barcodes. We then rounded the genotypes to remove the uncertainty during the comparison. Wilcoxon signed test p-value is indicated above the violin plots. B) Narrow-sense heritability measured with scRNA-seq consensus genotypes and non-multiplexed DNA sequencing. Three datasets are compared through a t-test on the bootstrap narrow-sense heritabilities. The consensus genotypes of the batch 1 F2 segregants obtained from the associated scRNA-seq data are represented in blue, the batch 1 F2 segregant genotypes in the reference dataset (bulk sequencing) are represented in yellow and the whole reference dataset of F2 segregants is represented in grey. For each dataset, the narrow-sense heritability has been measured from 500 random subsamples of 1000 F2 segregants. The bars represent the mean narrow-sense heritability, and the error bar length represents the 95% CI. The results are illustrated for the 15 environmental conditions where a minority of batch 1 F2 segregants have missing fitness data. The 23C-30C represents the temperature for the competition assay in YPD media while the other phenotypes represent growth on YNB, molasses (mol), mannose (Mann) or raffinose (raff) and chemical resistance to copper sulfate (Cu), ethanol (eth), guanidinium chloride (gu), lithium acetate (Li), Sodium dodecyl sulfate (SDS) and suloctidil (suloc) (29).

To further establish that the genotyping obtained from scRNA-seq data was comparable to previous non-multiplexed genotyping of the reference genotype panel, we estimated the contribution of genetic variation to the phenotypic variation, i.e. fitness heritability. Nguyen Ba and collaborators (2022) estimated the narrow- and broad-sense heritabilities of complex phenotypes associated with temperature gradient, carbon source and chemical resistance for which RM and BY segregants exhibit a significant level of diversity (29). We used our lineage assignment to that panel to obtain fitness but used our single-cell genotyping to perform this association (**Methods**). Encouragingly, all GCTA-REML estimates of narrow-sense heritability in the

A



B

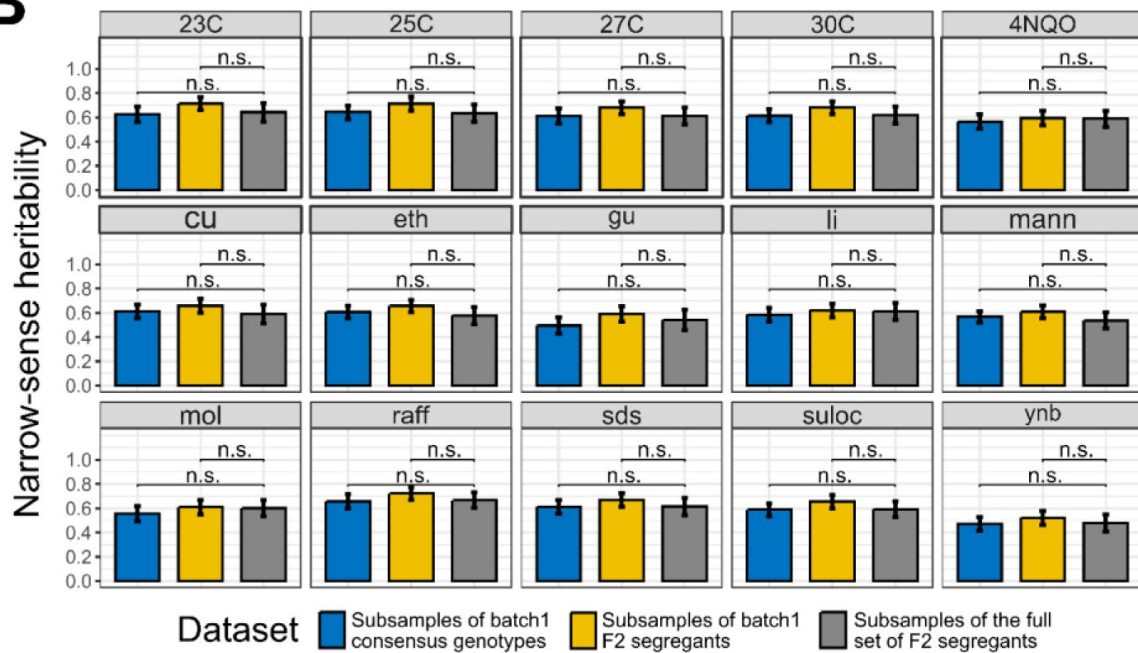


Figure 2

Single-cell RNA-seq data recapitulate bulk DNA and RNA assays results.

A) Effect of the HMM on the relatedness between single-cell genotypes and their closest reference lineage.

scRNA-seq batch 1 dataset (blue in [Figure 2B](#)) are similar to the estimates from Nguyen Ba and collaborators (2022) whether or not we focus on batch 1 (yellow and grey rectangles in [Figure 2B](#)).

Although the variance partitioning is consistent with previous studies, it only provides a broad view of the genotype-phenotype map as it does not allow to identify the loci that significantly explain phenotype variation. If the genotypes obtained by scRNA-seq were of high quality, then we would expect that a QTL mapping model from scRNA-seq would yield a similar model to non-multiplexed DNA sequencing data. To achieve this, we used a cross-validated stepwise forward linear regression on the strain fitness and consensus genotypes data from single cells that shared the same lineage assignment (**Methods**). Performing the QTL mapping on the batch 1 scRNA-seq dataset enabled the identification of 29 QTL compared to 31 QTL identified with the bulk barcoded approach (**Tables S2 and S3**). These QTL were largely similar as shown by the non-significant difference between the effect sizes (Wilcoxon signed rank test $p = 0.29$) and by a model similarity metric that considers the recombination distance between matched QTL, the similarity of the effect sizes and the allele frequencies (**Methods**). Using this approach, we estimated that the similarity score between the batch 1 single cells QTL and the batch 1 BB-QTL is 86.2% while each model respectively had a similarity score of 78.7% and 78.2% with the full BB-QTL mapping performed on the 99,950 F2 segregants (Figure S4). The QTL identified from the scRNA-seq dataset also recapitulated several important biological features of the reference panel such as an enrichment of non-synonymous and disordered region QTL (Figure S5).

Integrating scRNA-seq data to an existing GPM highlights selection on the transcriptome

Having shown that our scRNA-seq allows us to relate the cells to the corresponding F2 segregant fitness data via the genotype relatedness, we next sought to highlight new associations within the BY/RM GPM. Selection is often highlighted at the genotype level through convergent evolution, an increase in allele frequency within a population, or population genetics metric (34–36). However, the central dogma of molecular biology and the evolution tinkering model both entail that phenotype variation should be linked to transcriptomic variation. As our dataset included all these variables, we sought to evaluate the expression-genotype association strength and provide a variance partitioning framework to evaluate the association between the transcriptome and trait variation (**Methods**). The GCTA-REML variance partitioning model used to estimate trait heritability in the previous section can also be modified to include gene expression as the response variable and cell genotypes as the only random effect (**Methods**). This enables the quantification of expression heritability, i.e. the variance of expression explained by genotype. Using this approach on our large-scale dataset, we estimated that genotype explains 75.7% of expression variance, which is slightly higher than previous estimates from non-multiplexed eQTL mapping studies. Indeed, Albert and collaborators (2018) estimated that genotype explains 70% of expression variance using a lower-power dataset of 5,720 genes in 1,012 yeast segregants generated by the same parental strains (RM and BY).

In addition to expression heritability, the multidimensionality of our scRNA-seq dataset enables the evaluation of the association between traits and expression, which we demonstrate with the 30C phenotype (Figure 3).

The components of this variance partitioning all relate to at least one biological phenomenon. Indeed, the portion of trait variation explained exclusively by the genotype variation (red in Figure 3) represents the effect of mutations on fitness via several biological phenomena such as protein stability, enzymatic function etc, independent of expression level. For the 30C phenotype, this component explains 31.1% of the fitness variation in the BY/RM background which is slightly lower than the 36.8% explained by the shared component between phenotype, genotype and expression variations (purple in Figure 3). The latter represents the association between

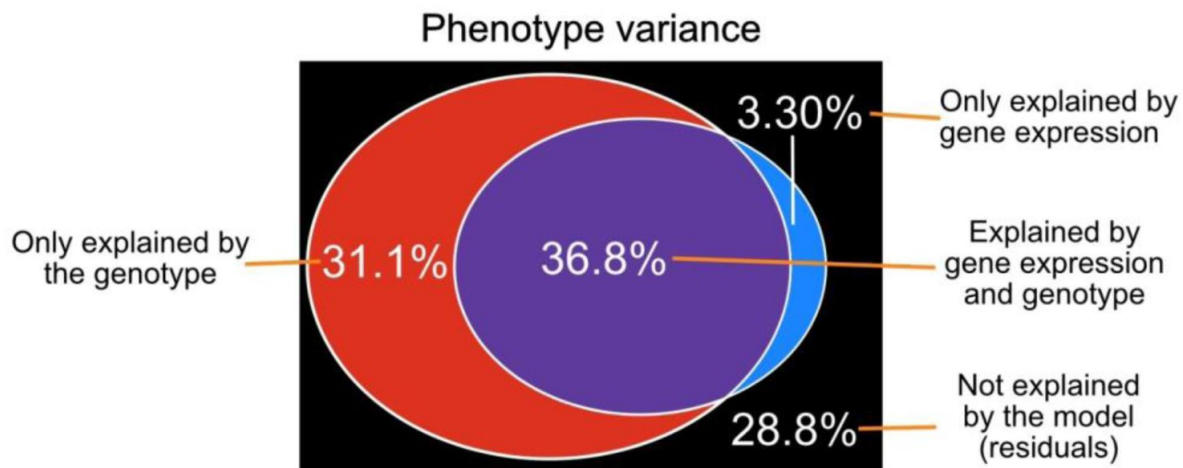


Figure 3

Variance partitioning of the 30C phenotype from scRNA-seq data.

The percentages represent the proportion of fitness variance (whole rectangle area) explained by the components and their overlaps. The ellipse area represents the phenotype variance explained by genotype or expression variation. The black area of the rectangle represents the residual of the model while the other colored areas represent the shared and exclusive components explaining fitness variation. The size of the shared component area (purple) is obtained by subtracting the variance explained by the model in [Equation 4](#) to the one in [Equation 2](#) or in [Equation 3](#) as it is a shared component of trait variation explained by genotype and expression variation. The size of the red genotype-exclusive component is obtained by subtracting the size of the shared component from the variance explained by the model in [Equation 2](#). The size of the blue expression-exclusive component is obtained by subtracting the size of the shared component from the variance explained by the model in [Equation 3](#). Finally, the size of the residuals is obtained by subtracting 1 by the sum of the size of all the defined components (red + blue + purple).

selection (fitness) and the transcriptome either through loci influencing fitness via expression directly or through loci affecting expression via an effect on cell fitness (indirectly) (37 [↗](#), 38 [↗](#)). Its considerable association with fitness variation thus supports the evolution tinkering model. As for the phenotype variation explained exclusively by gene expression (blue in **Figure 3** [↗](#)), it could represent epigenetics and stochastic gene expression, which weakly explain variations in the 30C phenotype.

Although this model accurately estimates the narrow-sense heritability of 30C, the residuals still represent 28.8% of fitness variation. This could be explained by unmeasured factors like high-order epistasis. However, the broad-sense heritability of this phenotype is similar to the narrow-sense heritability, suggesting that the residuals are mostly not explained by genotype and expression (29 [↗](#)). Nguyen Ba et al. (2022) also estimated that epistasis only explained around 5% of fitness (29 [↗](#)). Other factors like mitochondrial mutations, protein-protein interactions, or other protein properties could play a role in explaining the residual of the model and are good candidates to extend this analysis. These results suggest that a single run of scRNA-seq on a single batch of F2 yeast segregants converges with bulk DNA sequencing results while revealing previously hidden components of the GPM.

Revealing hidden components of the yeast GPM with scRNA-seq

Our integrative scRNA-seq approach is not limited to enabling the quantification of the association between transcriptomic changes and trait variation. Indeed, the same approach we used to identify QTL can be used to detect loci regulating gene expression which can reveal the cell mechanisms underlying trait variation through transcriptomic changes. We thus modified the QTL mapping framework such that the response variable is the level of expression of a single gene in the single cells (**Methods**). Because our cells are not synchronized, differences in the cell cycle could create transcriptomic variation that could obscure the association with genotypic variation. To control for this, we averaged out the effect of cell cycle by reconstructing the consensus expression and genotype profiles of each segregant significantly associated to the cells. This approach is a cost-efficient way to perform eQTL mapping from the expression profile and genotype of cells from thousands of lineages in a multiplexed way (sc-eQTL mapping).

Consistent with yeast non-multiplexed eQTL results, the genes with the highest expression heritability are enriched in functions related to carbohydrate catabolic process (GO:0016052) and cellular biosynthetic process (translation GO:0006412, organelle assembly GO:0070925, ribosome biogenesis GO:0042254 and gene expression GO:0010467) (Fisher's exact test FDR<0.05; **Methods**). In both datasets, these genes are also highly expressed, which reflects the positive correlation between expression heritability and expression levels ($R^2 = 0.66$ and $p < 2.2 \times 10^{-16}$). Although some of these highly expressed genes have important functions modulated by mutations, their higher heritability could also be explained by higher statistical power to measure heritability when expression counts are higher. Moreover, genes with the lowest expression heritability observed in the RM/BY background, which we defined as the bottom 10% expression heritability, are enriched in functions related to the cell cycle biological process (GO:0007049, Fisher's exact test FDR<0.05) which is consistent with bulk RNA-seq eQTL (21 [↗](#), 39 [↗](#)). This does not suggest that genes with this function are not important for variation in gene expression at the whole-genome scale but rather that mutations observed in the RM/BY background are not correlated with variation in the expression of these specific genes. This could reflect conservation for genes involved in the cell cycle biological process (40 [↗](#), 41 [↗](#)).

Because of the increased scale of our collection, our approach is more powered to estimate the gene expression heritability. Indeed, we detected a median of 21 eQTL per gene (mean = 34.3, s.d. = 29.7) which is almost 4 times higher than what has been detected by previous non-multiplexed RNA-seq (21 [↗](#)). We were also able to detect new overrepresented biological processes, i.e. DNA

metabolic process (GO:0006259) and the response to nutrient levels (GO:0031667), for which the variation of expression levels is weakly associated to the genetic variation observed across the F2 RM/BY segregants.

The functional enrichment analysis using scRNA-seq data revealed new associations between expression heritability and biological processes in the RM/BY genetic background. However, while it suggests that many eQTL are also QTL, it cannot accurately point to the specific loci involved in trait variation and cannot address whether mutations on regulatory hubs have stronger effects on traits. To investigate this, we mapped the QTL to hotspots of gene regulation (or regulatory hubs), which we defined as 25 kb genomic windows that were repeatedly identified in the eQTL mapping procedure (for different genes). This was done to acknowledge the uncertainty in the exact position of the eQTL due to linkage disequilibrium and power. We then ranked the 30C QTL identified by Nguyen Ba and collaborators (2022) based on their absolute effect size and correlated it to the rank of the eQTL hotspots based on the number of regulated genes. This resulted in a positive correlation (Spearman $\rho = 0.33$ and $p = 5.21 \times 10^{-5}$), suggesting that larger effects on the regulatory network translate into larger trait variation. Indeed, we observed that some previously reported high-effect-size QTL genes are located in eQTL hotspots, e.g. MKT1, HAP1, and IRA2 (Figure 4A [↗](#)).

Performing this rank test on individual genes also yielded the result that eQTL effect is correlated with the fitness effect for 35.1% of the genes (permutation test $p < 0.05$, see **Methods**). Although this correlation does not apply to most genes, it reveals potential regulatory mechanisms explaining the importance of the strongest growth loci or QTL. For instance, MKT1, i.e. the strongest growth loci, is part of a regulation hotspot affecting genes that are important for yeast growth like ENP1 which is involved in RNA processing and HXT6 which is involved in glucose uptake ([42](#) [↗](#)–[44](#) [↗](#)). Among the strongest growth loci, VPS70 is part of a hotspot of regulation that strongly affects the expression of RSF2, a zinc-finger protein regulating glycerol-based growth and respiration ([45](#) [↗](#)). Furthermore, the highest peak for expression regulation contains important growth loci in chromosome IV around the mating type loci. This suggests the presence of cells with different mating types in the dataset which we confirmed from the read mapping to Mat-a and Mat- α genes. This is consistent with previous budding yeast eQTL mapping and is also expected because the mating types in yeast express sets of genes that are “turned off” in other mating types ([20](#) [↗](#), [21](#) [↗](#), [46](#) [↗](#)). This peak of expression regulation is also responsible for regulating TDH3 which is involved in glycolysis and glucogenesis and can have an important effect on fitness ([47](#) [↗](#)).

These hotspots suggest that expression differences in BY/RM would predominantly be due to mutations in trans-regulatory elements. To test this, we partitioned the variation in gene expression between cis- and trans-regulatory loci for each gene (see **Methods**). This analysis revealed that all the genes are affected by at least one polymorphic trans-regulatory locus and that these polymorphic trans-regulatory loci explain most of that gene’s expression (Figure 4B [↗](#)). It is well known that mutations in promoters and nearby enhancers can influence gene expression ([48](#) [↗](#), [49](#) [↗](#)). Indeed, we identified many genes that contained an allele in a cis-regulatory element that strongly explain that gene’s expression variation ($n=750$ genes out of 6088, Figure 4B [↗](#)). As expected, mutations in cis-regulatory elements were of stronger effect size than trans-eQTL individually, but the cumulative aggregate effect of all trans-eQTL acting on that gene was comparable to the few cis-eQTL they had (Figure 4C [↗](#)). This can be explained by the fact that there are more opportunities for mutations to arise in trans-regulatory elements. Finally, we found that trans-eQTL have two times higher odds of affecting cell fitness than cis-eQTL ($\chi^2 p = 0.01$).

Furthermore, by repeating our eQTL hotspot analysis with Albert et al. (2018) data, we observed a non-significant association between eQTL hotspot and QTL ($\chi^2 p = 0.50$). In that bulk RNA-seq dataset, QTL have less odds of being in eQTL hotspots compared to other loci (od ratio = 0.72) while in this scRNA-seq dataset, QTL have twice more odds of being in eQTL hotspots. This could be

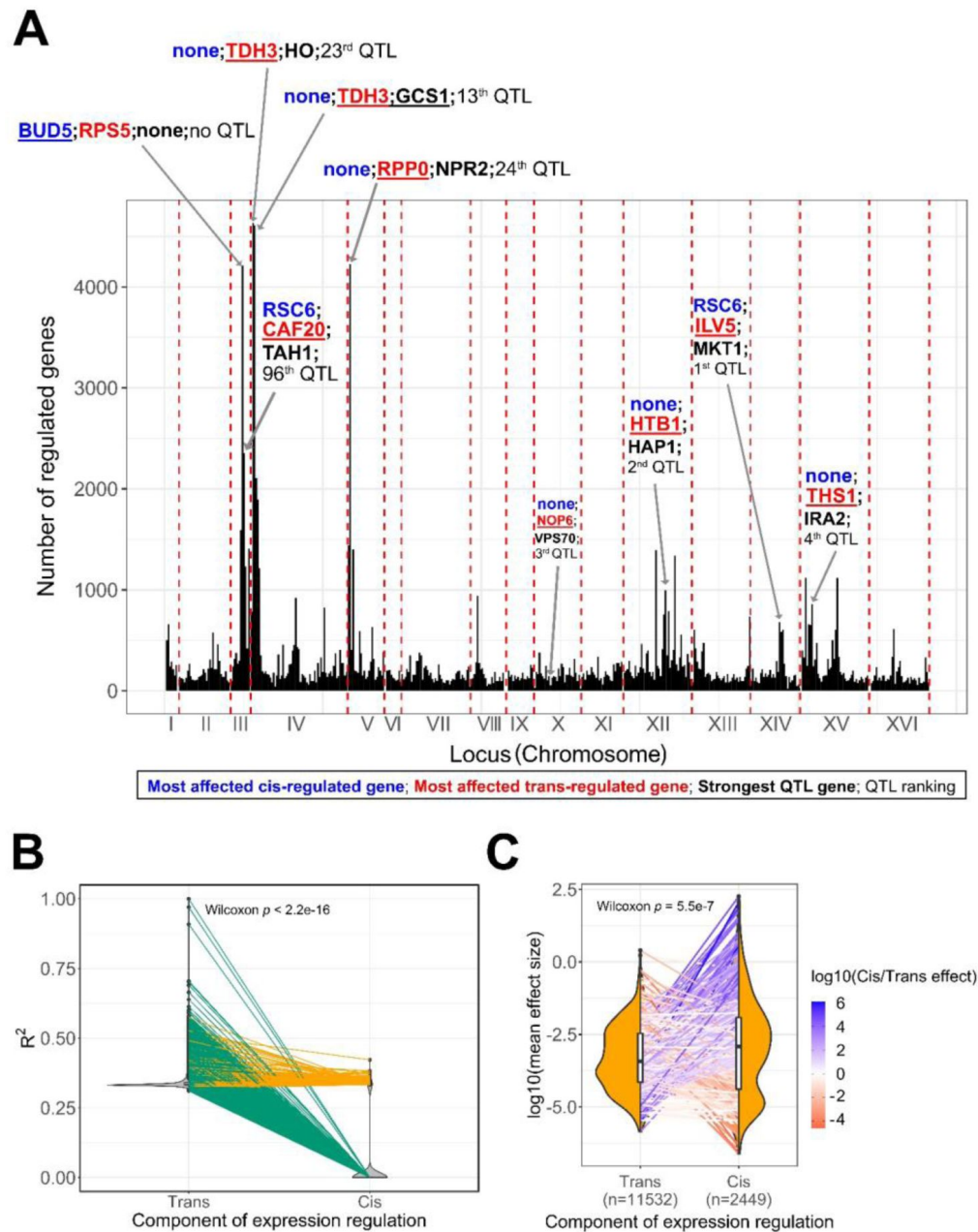


Figure 4

eQTL features underlying trait variation across the BY/RM segregants.

A) Mapping of the 30C QTL in the eQTL hotspots. We represent the hotspots of expression regulation as genomic windows (25 kb) to acknowledge the uncertainty around the real position of the eQTL due to linkage disequilibrium. We annotated the 5 top eQTL hotspots and the eQTL hotspots in which the top additive QTL identified by the BB-QTL mapping of the 30C phenotype are located. In these regions, we represented the most affected trans-regulated genes in red, the most affect cis-regulated gene in blue and the genes of the top QTL in black. The double quotation characters represent the absence of such genes in the associated region. We also represented the rank of the QTL in the set of 159 QTL of the 30C phenotype. B) Partitioning of the expression heritability or explained variance (R^2) among cis- and trans-eQTL. Each pair of points connected by a line represents a gene. Green lines represent the genes that are only have trans-eQTL and orange lines represent the genes that have both trans- and cis-eQTL. C) Comparison of the mean effect size between cis- and trans-eQTL. Each pair of points connected by a line represents a gene. The ratio of the average effect size between cis- and trans-eQTL is represented by the line color. The sample size of each eQTL category is represented in the x-axis. This is the number of trans-eQTL and cis-eQTL used for calculating the average effect sizes per gene not the number of points per distribution.

explained by the smaller sample size of that bulk RNA-seq dataset and by the fact that they control for differences in growth across lineages. Taken together, these results suggest that the link between the genetic basis of transcription variation across RM/BY segregants and fitness could only be revealed by integrating large-scale transcriptomic data to an existing GPM, which scRNA-seq facilitates.

Conclusion

By leveraging the scalability of scRNA-seq, we obtained thousands of transcriptomes from a reference pool of strains in a single experiment. This enabled the analysis of the association between genotype, transcriptome, and phenotype at an unprecedented scale. Questions surrounding transcriptomic variation and phenotypic variation have been at the center of many previous quantitative genetics studies (19–21,27,46,50–53). For instance, how strong is the connection between the transcriptome and trait variation? Although most GWAS loci are located at noncoding regions and are well-characterized for cis-regulation effects, is trans-regulation more important for phenotypic and expression variation? What are the gene functions that are the most heritable and modulable at the expression level? These ideas and discoveries all support the fact that researchers can gain valuable insight into the evolution of traits by integrating the transcriptome in GPM analyses, which can translate into fundamental knowledge or other important applications where phenotypes evolve.

In this study, we took advantage of a previously characterized BY/RM cross where the genetic basis of growth in various environments was examined in detail (29). By integrating transcriptomic data in this genotype-phenotype map, we revealed how transcriptomic components are involved in trait variation at a broad scale. Similar to a previous study, which obtained transcriptomes by individual strain sequencing, we found that gene expression is highly heritable. Further, our study design also allowed us to conclude that gene expression contributes to a significant portion of the phenotypic variation in this strain collection.

This finding is corroborated by our findings that most eQTL detected in our study were previously shown to be QTL. This is perhaps not surprising given that QTL in this cross were previously inferred to be in regulatory genes, but this provides a more mechanistic view of the effect of an allele on phenotype. Indeed, we find a bias for trans-regulation for generating transcription innovation where the cumulative effect of trans-eQTL on gene expression are significant. That is not to say that cis-regulatory alleles are dispensable as cis-regulatory alleles often have a large effect on gene expression. This is similar to studies on human transcriptomes where trans-eQTL are cumulatively more impactful and enriched in sets of complex trait loci (19,51–53). This genome-wide view of the genetic basis of transcriptional variation has consequences for the evolution of phenotypes, as the target size afforded by trans-eQTL is far larger than cis-eQTL. Thus, adaptation to small and fluctuating environmental changes may proceed preferentially through allelic changes or recombination of many small-effect trans-eQTL, but large expression changes are likely to require some cis-eQTL.

In this study, we leveraged the availability of genotype and phenotype data on our pool of strains. This was obtained by liquid handling robotics and pooled competitive growth assay with barcode sequencing. While this was performed on a very large scale, it was essentially obtained by brute-force and through approaches that are not necessarily applicable to other systems. While it is clear from our results that single-cell sequencing can achieve the same genotype quality as single-reaction genotyping, it is much harder to obtain phenotype data from scRNA-seq. Thus, our framework might not be readily translatable to other systems where similar studies on the GPM are desirable. However, two observations from this cross can be used to suggest an experimental approach. First, while epistasis is important, it contributes to a relatively small portion of the phenotypic variance. Second, transcriptomic variation contributes little to the missing heritability.

Thus, it may be possible to use predicted fitness instead of observed fitness and recapitulate essentially similar results as this study. Indeed, the phenotype under study has a negligible component of its variance explained by transcription independent of the genotype, but this feature was unknown before the study. This alternative approach can only relate expression to the heritable component of traits, which would have missed this variance explained by transcription independent of the genotype and explains why we did not employ it. Predicted fitness could be obtained from bulk-segregant analysis where the additive effect of loci can be inferred from whole-genome sequencing (28 [4](#),54 [4](#)). However, it is not clear if the proposed experimental approach is generalizable in any other context.

However, despite the study's limitation on generalizability, our scRNA-seq framework helps bridge the understanding of how genetic variation influences transcriptomic variation at a broad scale. Our framework relies on identifying the genome of single cells from the transcriptome, which is going to be possible from low-coverage sequencing when genetic variation within the pool is high (such as this cross, microbiome sequencing, or cancer cells with extensive copy number variation), and from low cell diversity with sufficient transcriptomic variation such that aggregation of single-cells with similar transcriptomes can afford pseudo-high coverage sequencing. Although we restricted the scope of this work to broad-scale transcriptomic and genomic patterns related to trait variation, we demonstrated that it is possible to integrate large-scale transcriptome data into an existing GPM, which can be exploited in the future for more fine-scale analysis like co-expression analysis and highlighting how specific trans-regulatory patterns within the cell network lead to trait variation. Thus, integrating genotype, transcriptome, and phenotype using scRNA-seq data can be particularly efficient for developing a better understanding of other important traits or diseases.

Material and methods

Yeast strains and segregants

We analyzed cells from a single batch (batch 1) of 4,489 F2 segregants obtained from a F1 cross between the yeast laboratory strain BY4741 and the vineyard strain RM11-1a generated in a previous study (29 [4](#)). These strains have been selected to generate this collection of segregants because they exhibit differences in multiple phenotypes including the adaptation to temperature, the ability to process different sources of carbon and the ability to resist antifungal compounds. Therefore, the genetic variation observed across the segregants can be correlated to the differences in growth rate observed in the 18 environments recapitulating these phenotypes in the Nguyen et al (2022) study (29 [4](#)). The selection of the batch is random and the fact that we performed the analyses on a single batch eliminates batch effects that could obscure variable associations. Genotypes and fitness data used were the same ones obtained in the previous study.

Yeast growth and single-cell RNA-sequencing protocol

To prepare strains for scRNA sequencing, we unfroze the batch of segregants and inoculated approximately 5×10^6 cells in YPD (1% Yeast Extract, 2% Peptone, 2% Dextrose) to saturation. The next day, about 10^7 cells were passaged to 5 mL of fresh YPD and grew for 4 hours to bring cells to log-phase. We then pelleted 100 μ L of cells and resuspended them in spheroplasting solution (5 mg/mL zymolyase 20T, 10 mM DTT, 1 M Sorbitol, 100 mM Sodium Phosphate pH 7.4) at a concentration of 10^7 cells/mL. The cells were incubated at 37 degrees Celsius for approximately 10 minutes at which point spheroplasting was verified by mixing a small aliquot of cells with detergent to observe lysis. The cells at this point were quantified using a hemocytometer and prepared using the standard 10x Genomics Gel Beads-in-emulsion (GEM) protocol. We used the Chromium Next GEM Single-cell 3' Reagent Kit to prepare the sequencing libraries and sequenced on a NextSeq 500 high-output flow cell.

We note that the cells analyzed here were grown in bulk and assayed for their transcriptome in log-phase. Our fitness data was obtained from competitive bulk fitness assays which include several whole growth cycles over multiple days and thus capture lag phase, exponential growth, and saturation. Nevertheless, previous experiments had shown that fitness was mostly determined by exponential growth which suggests that our analysis is adequate even if the cells were prepared for sequencing at a single time point.

Single-cell RNA-sequencing data parsing

From the scRNA-seq reads, we obtained gene expression levels and allele counts using the pipeline count from CellRanger version 3.1.0 (55). For each of the ancestral strains, i.e RM11-1a and BY4741, the pipeline mapped the scRNA-seq reads to the reference genome, filtered the barcodes by comparing the UMI count per barcode distribution to a background model of empty gel-bead in-emulsion, and counted the number of UMI per gene per barcode. The barcode filtering retained 18,233 barcodes. For each barcode, we then counted the number of RM and BY alleles at each polymorphic site by parsing the RM and BY bam files using a Python script (https://github.com/arnaud00013/sc-eQTL/tree/main/II_scRNA-seq_genotyping). This script only keeps reads that mapped at the same loci on both reference genomes to increase the level of confidence of the mapping.

Correction and imputation of single-cell genotypes with a Hidden Markov Model

Because there are only two possible alleles at each polymorphic sites of the F2 RM/BY segregants, their genotype can be recapitulated by a quantitative variable measuring the proportion of reads from one of the parental strains, which is RM in our dataset. The raw allele count data provides a first estimate of this RM allele frequency at each polymorphic site. However, due to the low mean depth of coverage of scRNA-seq data (0.2x), the absence of reads in some polymorphic sites and the biases introduced during sequencing like index hopping/swapping, we expect that the raw data can be imputed and corrected for errors and uncertainty in the observed alleles. Therefore, we applied a Hidden Markov Model (HMM) on the observed allele count. Such model can infer accurate genotype data at sequencing depths as low as 0.1x (29,31,56). Nguyen Ba and collaborators (2022) designed an HMM to infer the segregants genotypes from bulk DNA sequencing by accounting for sequencing error rate, recombination rate and index swapping rate (29). Because scRNA-seq uses the reverse transcriptase, which has a higher error rate, and because it is a pooled assay with higher chances of index swapping, we expected the HMM parameter to differ for the single cell data. Therefore, we adapted the HMM to scRNA-seq data by measuring its parameters in our dataset (Figure S1). The scripts are available on GitHub (https://github.com/arnaud00013/sc-eQTL/tree/main/II_scRNA-seq_genotyping).

Assigning single cells to the reference panel strains

To evaluate the level of relatedness between the reference panel strains and the imputed single cell genotypes, we used the expected distance to identify the strain that best relate to each single cell:

$$\text{Expected distance}(g_c, g_s) = \sum_{i=1}^{41594} g_c + g_s - 2g_c g_s \quad (\text{Eq.1})$$

where g_c is the cell genotype and g_s is the strain genotype. Next, we assigned the single cell to its best match in the studied batch of 4,489 strains only if this match is better than the best match in randomly generated batches of the same size (Figure S2). This procedure is implemented and available at https://github.com/arnaud00013/sc-eQTL/tree/main/III_Genotype_analysis.

Partitioning the phenotypic variance into genetic and transcriptomic components

To analyze the yeast GPM at a broad scale and to evaluate the association between selection and the transcriptome, we estimated the contribution of genetic and transcriptomic variations to phenotypic variation from scRNA-seq data. More precisely, we performed a Genome-wide Complex Trait Analysis (GCTA) by fitting a linear mixed model to the data using the restricted maximum-likelihood (REML) method (57 [↗](#)):

$$y = X\beta + W_g u_g + \varepsilon_g \quad (\text{Eq.2})$$

$$y = X\beta + W_e u_e + \varepsilon_e \quad (\text{Eq.3})$$

$$y = X\beta + W_g u_g + W_e u_e + \varepsilon \quad (\text{Eq.4})$$

where y is the fitness vector for the n segregants with a significant association to batch 1 cells, X is the $n \times k$ matrix of k fixed effects, β is the vector of k coefficients of the fixed effects, W_g is the $n \times p$ genotype matrix, u_g is the vector of p SNP effects, W_e is the $n \times m$ expression matrix, u_e is the vector of m gene expression effects and ε is the error term. To decrease the effect of gene dropouts and noise in the gene expression profile of each cell, we fitted an imputation model called DISCERN to the expression data (33 [↗](#)). DISCERN is an autoencoder that learns how to reconstruct the expression profile of high-coverage cells after embedding it in a lower dimension. These operations eliminate the gene dropouts and denoises the expression profile. We also controlled for the effect of the cell cycle on gene expression by using the consensus F2 segregants expression and genotypes from the associated cell data. The consensus data has been obtained using the cell with the maximum read count. Because the dataset does not include fixed effects, we set the fixed effect to a vector of ones such that its coefficients represent the mean fitness while the genotype and expression data are the random effects that explain the fitness variance along with the error terms. The REML solution assumes that the data follow a Gaussian distribution, so the data are standardized before fitting the model. We also divided the standardized expression counts by the cell sum of expression counts to control for molecule count biases across cells. The cell fitness is based on the fitness of the closest segregant in batch 1 as measured by the expected distance. Because this model is linear and additive, it can be compared to the estimates of narrow-sense heritability obtained by Nguyen Ba and collaborators (2022) (29 [↗](#)). The difference between the variance explained in equation 4 [↗](#) and equations 2 [↗](#) or 3 [↗](#) allow to infer the variance explained only by the genotype or the expression component of the model. The code for the variance partitioning is available on GitHub (https://github.com/arnaud00013/sc-eQTL/tree/main/IV_variance_partitioning [↗](#)).

Estimating the expression heritability from scRNA-seq

To obtain this estimate from scRNA-seq data, we considered the fact that GCTA-REML only takes a vector as a response variable while the gene expression matrix is multi-dimensional. To solve this, we orthogonalized the gene expression matrix using principal component analysis (PCA), and used its first PC, which explains 99.6% of the variance, as a response variable of the model. In case 99% of the expression matrix is explained by more than one PC, it is possible to use a weighted sum to obtain the expression heritability. Indeed, if the expression PCs recapitulate the total expression variance and are orthogonal or independent to each other, then the sum of the PCs variance explained by genotype should be the expression heritability. If " k " is the number of PCs explaining at least 99% of expression variance:

$$\text{Expression heritability} = \sum_{i=1}^k PC_i \text{ eigen value} * "PC_i \sim \text{genotype}" \text{ model } R^2 \quad (\text{Eq.5})$$

QTL mapping

To identify the loci that influence cell fitness, we performed a linear regression on the consensus genotypes of the strains from the scRNA-seq data and the strain fitness. We decided to use the consensus genotypes of the strains as they relate better to the bulk segregant genomes. To build the consensus genotypes, we defined cells from the same lineage as the ones that shared the same closest F2 segregant in batch 1. Next, we used the median to obtain cells' consensus genotypes as it is less sensitive to outliers and because it yields the best relatedness to the batch 1 reference genotypes (median $R^2 = 87.0\%$; $\mu=79.5\%$; $\sigma=18.2$; **Figure S3**). We selected the QTL in the linear models using cross-validation on the scRNA-seq data. This analysis consists in dividing the dataset into 10 random partitions of similar sample sizes and running a cross-validated stepwise forward linear regression on each partition. For each partition, the model starts with no QTL and a linear model "Fitness ~ Genotype" is fitted using the genotype data at each polymorphic site, where the correlation coefficient represents the effect size of the SNP. Then, the forward search starts and at each iteration, a new locus with the minimum linear model residual sum of squares (RSS) is added to the QTL model, which is updated with new effect sizes after the addition of a new SNP. Because the order of addition of QTL matters in the forward search and because some QTL are linked or collinear, the model can be refined by exploring different QTL around the local optima. These steps are repeated until the model RSS cannot be improved anymore or until the number of QTL reaches an arbitrary maximum far from the cross-validated number of QTL. After the forward search is completed in each partition, the algorithm calculates the optimal λ values that minimizes the objective function F_o :

$$F_o(\beta) = RSS(\beta) + Lasso\ penalty(\beta)$$

$$\|Y - X\beta\|_2^2 + \lambda\|\beta\|_0 \text{ (Eq.6)}$$

where β is the vector of SNP effect sizes in the QTL model, $\|Y - X\beta\|_2^2$ is the RSS of the linear QTL model, λ defines the penalty for adding a new SNP to the model and $\|\beta\|_0$ is the number of SNPs in the QTL model. This objective function has the property to add sparsity in the QTL model and thus avoid overestimating the number of QTL while being consistent (29). The optimal λ has a minimum of $\log(n)$ which corresponds to the Bayesian Information Criterion (BIC), which is known to yield correct models asymptotically (58). This allows to consider the possibility that a sparser model than the one found using the BIC could yield better predictive power on a test set while avoiding overfitting. The optimal λ values found in all the partitions are then averaged and the resulting mean λ is used to solve the objective function in the full dataset, which yields the optimal QTL model. The cross-validation assumes that the partitions are independent, such that the variance explained by the model and the number of relevant QTL are unbiased estimates.

Highlighting hotspots of gene regulation through eQTL mapping

To identify the loci regulating gene expression regulation, we adapted the QTL mapping framework using expression as the predicted phenotype. Because this approach had to be repeated for each of the 6,240 genes, we needed to modify it so that the execution time was convenient. To do so, the parameter λ was not estimated using cross-validation but rather from the Bayesian Inference Criterion (BIC), i.e. $\lambda = \log(n)$ where n is the number of cells. We found that the BIC was often selected by the cross-validation procedure when tested on a few genes and thus we do not believe that this approach will significantly change our results.

To acknowledge the uncertainty around the exact position of eQTL due to linkage disequilibrium, we define eQTL hotspots as 25 kb genomic windows that were repeatedly identified in the eQTL mapping procedure. Finally, we defined hotspots of gene regulation as genomic windows in the fourth quantile of the distribution of the number of regulated genes per window and the fourth quantile of the number of eQTL per window. The code for the single cell eQTL mapping is available on GitHub (https://github.com/arnaud00013/sc-eQTL/tree/main/V_sc_eQTL_mapping).

Functional enrichment analysis by gene ontology annotation

To highlight gene functions enriched at different levels of expression or expression heritability, we performed the panther database binomial test for statistical overrepresentation of gene ontology biological processes (39,59). A low level was defined as within the 25% bottom part of the distribution (<Q1) while a high level was defined as within the top 25% part of the distribution (>Q3). The *p*-values were corrected for multiple testing using the false discovery rate correction (FDR).

Matching QTL to eQTL

To evaluate the contribution of gene expression regulation to fitness variation, we created a model to match QTL and eQTL based on the similarity of loci and the similarity of predicted effect on gene expression. More precisely, for each of the 6,088 genes for which we could detect eQTL, we performed a new eQTL model by correlating the expression level of the gene to the genetic variation at QTL positions. This allowed us to measure the predicted effect of the QTL on gene expression. We then calculated the distance between the QTL and the real eQTL of the gene based on recombination distance within each chromosome, which decreases exponentially with genetic distance, and the difference in the predicted effect on the gene expression using the formulation developed by Nguyen Ba et al (2022) (29). Next, we used the same Needleman-Wunsch algorithm to find the most likely set of pairing between QTL and eQTL, where an unmatched QTL is also possible but penalized. Finally, we determined the proportion of genes for which gene expression regulation is associated with higher fitness. To do so, for each gene, we performed a permutation test by comparing the average rank of the matched QTL of the gene to the average rank of 999 random subsets of unmatched QTL of the same size. The *p*-value is the proportion of random subsets of unmatched QTL with a higher average QTL rank than the set of matched QTL.

Comparing cis- and trans-eQTL contribution to expression variation

We used the definition of local eQTL in Albert et al. (2018) to define cis-eQTL, i.e. any eQTL between 1,000 bp upstream of the gene and 200 bp downstream of the gene. Thus, we defined trans-eQTL as the eQTL that do not follow this criterion. For each gene, we then performed variance partitioning using the GCTA:

$$y = X\beta + W_{g_cis}u_{g_cis} + \varepsilon_{cis} \quad (\text{Eq.7})$$

$$y = X\beta + W_{g_trans}u_{g_trans} + \varepsilon_{trans} \quad (\text{Eq.8})$$

$$y = X\beta + W_{g_cis}u_{g_cis} + W_{g_trans}u_{g_trans} + \varepsilon \quad (\text{Eq.9})$$

where *y* is the vector of expression level of the gene across the *n* cells, *X* is the *n*×*k* matrix of *k* fixed effects, *β* is the vector of *k* coefficients of the fixed effects, *W_{g_cis}* is the *n*×*p* cis-eQTL genotype matrix, *u_{g_cis}* is the vector of *p* cis-eQTL effects on expression, *W_{g_trans}* is the *n*×*m* trans-eQTL expression matrix, *u_g* is the vector of *m* trans-eQTL effects on expression and *ε* represent the error terms. Because the dataset does not include fixed effects, we set the fixed effect to a vector of ones such that its coefficients represent the mean expression level while the cis-eQTL and trans-eQTL genotypes are the random effects that explain the expression variance along with the error terms. We can infer the variance explained by the cis-eQTL by the difference in variance explained between the models in equations 9 and 8. Likewise, the difference of variance explained by the models in equations 9 and 7 can help us estimate the variance explained by the trans-eQTL. Finally, we estimate the effect sizes using the absolute value of the correlation coefficients of each loci and compare the mean between the cis- and trans-eQTL from the same gene (paired data) with a Wilcoxon signed rank test.

Data availability

The code used for this study is available and explained at <https://github.com/arnaud00013/sc-eQTL> and the original single-cell reads from the pooled segregants scRNA-seq assay have been uploaded in the NCBI BioProject database with the accession number PRJNA1022775. The single-cell barcodes expression data are also available at <https://github.com/arnaud00013/sc-eQTL> as an archive file named `Matrix_gene_expression_barcodes_1_to_9000.csv.tar.gz` or `Matrix_gene_expression_barcodes_9001_to_18233.csv.tar.gz`.

Acknowledgements

ANNB acknowledges support from the Natural Science and Engineering Research Council of Canada (NSERC RGPIN-2021-02716 and DGEER-2021-00117) and AN acknowledges support from the Natural Science and Engineering Research Council (NSERC CGS-D App Id: 569340-2022, UTF Fellowship from the University of Toronto and Rustom H. Dasture Graduate Scholarship in Cell and Systems Biology). The computations in this paper were enabled by resources provided by Compute Canada. We also thank the Centre for the Analysis of Genome Evolution and Function (CAGEF) and TCAG at SickKids for sequencing support.

References

1. Bartoli C, Roux F (2017) **Genome-Wide Association Studies In Plant Pathosystems: Toward an Ecological Genomics Approach** *Front Plant Sci* **8** <https://doi.org/10.3389/fpls.2017.00763>
2. Ferreira MA, Gamazon ER, Al-Ejeh F, Aittomaki K, Andrulis IL, Anton-Culver H, et al. (2019) **Genome-wide association and transcriptome studies identify target genes and risk loci for breast cancer** *Nat Commun* **10** <https://doi.org/10.1038/s41467-018-08053-5>
3. Aguet F, Alasoo K, Li YI, Battle A, Im HK, Montgomery SB, et al. (2023) **Molecular quantitative trait loci** *Nat Rev Methods Primer* **3**:1–22
4. Tarantino LM, Eisener-Dorman AF (2012) **Forward Genetic Approaches to Understanding Complex Behaviors** *Curr Top Behav Neurosci* **12**:25–58
5. Casamassimi A, Federico A, Rienzo M, Esposito S, Ciccodicola A (2017) **Transcriptome Profiling in Human Diseases: New Advances and Perspectives** *Int J Mol Sci* **18**
6. Wainberg M, Sinnott-Armstrong N, Mancuso N, Barbeira AN, Knowles DA, Golan D, et al. (2019) **Opportunities and challenges for transcriptome-wide association studies** *Nat Genet* **51**:592–9
7. Adams J (2008) **Transcriptome: Connecting the Genome to Gene Function**
8. Williams CG, Lee HJ, Asatsuma T, Vento-Tormo R, Haque A (2022) **An introduction to spatial transcriptomics for biomedical research** *Genome Med* **14**
9. King MC, Wilson AC (1975) **Evolution at Two Levels in Humans and Chimpanzees** *Science*
10. Jacob F (1977) **Evolution and Tinkering** *Science*
11. Hoekstra HE, Coyne JA (2007) **The Locus of Evolution: Evo Devo and the Genetics of Adaptation** *Evolution* **61**:995–1016
12. Kratochwil CF, Meyer A. (2015) **Evolution: Tinkering within Gene Regulatory Landscapes** *Curr Biol* **25**:R285–8
13. Primig M, Williams RM, Winzeler EA, Tevzadze GG, Conway AR, Hwang SY, et al. (2000) **The core meiotic transcriptome in budding yeasts** *Nat Genet* **26**:415–23
14. Cavalieri D, Townsend JP, Hartl DL (2000) **Manifold anomalies in gene expression in a vineyard isolate of *Saccharomyces cerevisiae* revealed by DNA microarray analysis** *Proc Natl Acad Sci U S A* **97**:12369–74
15. Maurano MT, Humbert R, Rynes E, Thurman RE, Haugen E, Wang H, et al. (2012) **Systematic Localization of Common Disease-Associated Variation in Regulatory DNA** *Science* **337**:1190–5
16. Schaub MA, Boyle AP, Kundaje A, Batzoglou S, Snyder M (2012) **Linking disease associations with regulatory information in the human genome** *Genome Res* **22**:1748–59

17. Nicolae DL, Gamazon E, Zhang W, Duan S, Dolan ME, Cox NJ (2010) **Trait-Associated SNPs Are More Likely to Be eQTLs: Annotation to Enhance Discovery from GWAS** *PLoS Genet* **6**
18. Dutta D, He Y, Saha A, Arvanitis M, Battle A, Chatterjee N (2022) **Aggregative trans-eQTL analysis detects trait-specific target gene sets in whole blood** *Nat Commun* **13**
19. Wang L, Babushkin N, Liu Z, Liu X (2024) **Trans-eQTL mapping in gene sets identifies network effects of genetic variants** *Cell Genomics* **4**
20. Brem RB, Yvert G, Clinton R, Kruglyak L. (2002) **Genetic Dissection of Transcriptional Regulation in Budding Yeast**
21. Albert FW, Bloom JS, Siegel J, Day L, Kruglyak L., Wittkopp PJ (2018) **Genetics of trans-regulatory variation in gene expression** *eLife* **7**
22. Schwarz T, Boltz T, Hou K, Bot M, Duan C, Loohuis LO, et al. (2022) **Powerful eQTL mapping through low-coverage RNA sequencing** *Hum Genet Genomics Adv* **3**
23. Fan Y, Zhu H, Song Y, Peng Q, Zhou X (2020) **Efficient and effective control of confounding in eQTL mapping studies through joint differential expression and Mendelian randomization analyses** *Bioinformatics* **37**:296–302
24. Bush WS, Moore JH (2012) **Chapter 11: Genome-Wide Association Studies** *Plos Comput Biol* **8** <https://doi.org/10.1371/journal.pcbi.1002822>
25. Lorenz K, Cohen BA (2012) **Small- and Large-Effect Quantitative Trait Locus Interactions Underlie Variation in Yeast Sporulation Efficiency** *Genetics* **192**
26. Shaffer SM, Dunagin MC, Torborg SR, Torre EA, Emert B, Krepler C, et al. (2017) **Rare cell variability and drug-induced reprogramming as a mode of cancer drug resistance** *Nature* **546**
27. Bloom JS, Ehrenreich IM, Loo WT, Lite TLV, Kruglyak L (2013) **Finding the sources of missing heritability in a yeast cross** *Nature* **494**:234–7
28. Ehrenreich IM, Torabi N, Jia Y, Kent J, Martis S, Shapiro JA, et al. (2010) **Dissection of genetically complex traits with extremely large pools of yeast segregants** *Nature* **464**:1039–42
29. Ba AN Nguyen, Lawrence KR, Rego-Costa A, Gopalakrishnan S, Temko D, Michor F, et al., Verstrepn KJ (2022) **Barcoded Bulk QTL mapping reveals highly polygenic and epistatic architecture of complex traits in yeast** *eLife* **11**
30. Vermeersch L, Jariani A, Helsen J, Heineike BM, Verstrepn KJ., Devaux F (2022) **Single-Cell RNA Sequencing in YeastYeast Using the 10x Genomics Chromium Device** *Yeast Functional Genomics: Methods and Protocols* https://doi.org/10.1007/978-1-0716-2257-5_1
31. Hwang B, Lee JH, Bang D (2018) **Single-cell RNA sequencing technologies and bioinformatics pipelines** *Exp Mol Med* **50** <https://doi.org/10.1038/s12276-018-0071-8>
32. Jariani A, Vermeersch L, Cerulus B, Perez-Samper G, Voordeckers K, Van Brussel T, et al., Rokas A, Barkai N (2020) **A new protocol for single-cell RNA-seq reveals stochastic gene expression during lag phase in budding yeast** *eLife* **9**

33. Hausmann F, Ergen C, Khatri R, Marouf M, Hänzelmann S, Gagliani N, et al. (2023) **DISCERN: deep single-cell expression reconstruction for improved cell clustering and cell subtype and state detection** *Genome Biol* **24**
34. Bergström A, Simpson JT, Salinas F, Barré B, Parts L, Zia A, et al. (2014) **A High-Definition View of Functional Genetic Variation from Natural Yeast Genomes** *Mol Biol Evol* **31**:872–88
35. 35., Good BH, McDonald MJ, Barrick JE, Lenski RE, Desai MM. (2017) **The dynamics of molecular evolution over 60,000 generations** *Nature* **551**:45–50
36. Johnson MS *et al.* (2021) **Phenotypic and molecular evolution across 10,000 generations in laboratory budding yeast populations** *eLife*
37. Sun XM, Bowman A, Priestman M, Bertaux F, Martinez-Segura A, Tang W, et al. (2020) **Size-Dependent Increase in RNA Polymerase II Initiation Rates Mediates Gene Expression Scaling with Cell Size** *Curr Biol* **30**:1217–1230
38. Marguerat S, Bähler J (2012) **Coordinating genome expression with cell size** *Trends Genet* **28**:560–5
39. Mi H, Muruganujan A, Huang X, Ebert D, Mills C, Guo X, et al. (2019) **Protocol Update for large-scale genome and gene function analysis with the PANTHER classification system (v.14.0)** *Nat* **14**:703–21
40. Bähler J (2005) **Cell-cycle control of gene expression in budding and fission yeast** *Annu Rev Genet* **39**:69–94
41. Yu L, Castillo LP, Mnaimneh S, Hughes TR, Brown GW (2006) **A Survey of Essential Gene Function in the Yeast Cell Division Cycle** *Mol Biol Cell* **17**:4736–47
42. Roos J, Luz JM, Centoducati S, Sternglanz R, Lennarz WJ (1997) **ENP1, an essential gene encoding a nuclear protein that is highly conserved from yeast to humans** *Gene* **185**:137–46
43. Chen W, Bucaria J, Band DA, Sutton A, Sternglanz R (2003) **Enp1, a yeast protein associated with U3 and U14 snoRNAs, is required for pre-rRNA processing and 40S subunit synthesis** *Nucleic Acids Res* **31**:690–9
44. Roy A, Dement AD, Cho KH, Kim JH (2015) **Assessing Glucose Uptake through the Yeast Hexose Transporter 1 (Hxt1)** *PLOS ONE* **10**
45. Lu L, Roberts GG, Oszust C, Hudson AP (2005) **The YJR127C/ZMS1 gene product is involved in glycerol-based respiratory growth of the yeast *Saccharomyces cerevisiae*** *Curr Genet* **48**:235–46
46. Brem RB, Storey JD, Whittle J, Kruglyak L (2005) **Genetic interactions between polymorphisms that affect gene expression in yeast** *Nature* **436**:701–3
47. Zande P Vande, Hill MS, Wittkopp PJ (2022) **Pleiotropic effects of trans-regulatory mutations on fitness and gene expression** *Science* **377**:105–9
48. Mattioli K, Oliveros W, Gerhardinger C, Andergassen D, Maass PG, Rinn JL, et al. (2020) **Cis and trans effects differentially contribute to the evolution of promoters and enhancers** *Genome Biol* **21**

49. Romero IG, Ruvinsky I, Gilad Y (2012) **Comparative studies of gene expression and the evolution of gene regulation** *Nat Rev Genet* **13**:505–16
50. Bloom JS, Boocock J, Treusch S, Sadhu MJ, Day L, Oates-Barker H, et al., Landry CR, Barkai N (2019) **Rare variants contribute disproportionately to quantitative trait variation in yeast** *eLife* **8**
51. Westra HJ, Peters MJ, Esko T, Yaghootkar H, Schurmann C, Kettunen J, et al. (2013) **Systematic identification of trans-eQTLs as putative drivers of known disease associations** *Nat Genet* **45**:1238–43
52. Yao C, Joehanes R, Johnson AD, Huan T, Liu C, Freedman JE, et al. (2017) **Dynamic Role of trans Regulation of Gene Expression in Relation to Complex Traits** *Am J Hum Genet* **100**:571–80
53. Liu X, Li YI, Pritchard JK (2019) **Trans Effects on Gene Expression Can Drive Omnigenic Inheritance** *Cell* **177**:1022–1034
54. Brauer MJ, Christianson CM, Pai DA, Dunham MJ (2006) **Mapping novel traits by array-assisted bulk segregant analysis in *Saccharomyces cerevisiae*** *Genetics* **173**:1813–6
55. Zheng GX, Terry JM, Belgrader P, Ryvkin P, Bent ZW, Wilson R, et al. (2017) **Massively parallel digital transcriptional profiling of single cells** *Nat Commun* **8**
56. Stark R, Grzelak M, Hadfield J (2019) **RNA sequencing: the teenage years** *Nat Rev Genet* **20**:631–56
57. Yang J, Lee SH, Goddard ME, Visscher PM. (2011) **GCTA: a tool for genome-wide complex trait analysis** *Am J Hum Genet* **88**:76–82
58. Harrell FE. (2001) **Regression Modeling Strategies: With Applications to Linear Models, Logistic Regression, and Survival Analysis**
59. Thomas PD., Ebert D., Muruganujan A., Mushayahama T., Albou L.P., Mi H. (2023) **PANTHER: Making genome-scale phylogenetics accessible to all** *Protein Science*

Editors

Reviewing Editor

Vaughn Cooper

University of Pittsburgh, Pittsburgh, United States of America

Senior Editor

Detlef Weigel

Max Planck Institute for Biology Tübingen, Tübingen, Germany

Reviewer #1 (Public review):

In the revision of their paper, N'Guessan et al have improved the report of their study of expression QTL (eQTL) mapping in yeast using single cells. The authors make use of advances in single cell RNAseq (scRNAseq) in yeast to increase the efficiency with which this type of analysis can be undertaken. Building on prior research led by the senior author that entailed genotyping and fitness profiling of almost 100,000 cells derived from a cross between two

yeast strains (BY and RM) they performed scRNAseq on a subset of ~5% ($n = 4,489$) individual cells. To address the sparsity of genotype data in the expression profiling they used a Hidden Markov Model (HMM) to infer genotypes and then identify the most likely known lineage genotype from the original dataset. To address the relationship between variance in fitness and gene expression the authors partition the variance to investigate the sources of variation. They then perform eQTL mapping and study the relationship between eQTL and fitness QTL identified in the earlier study.

This paper seeks to address the question of how quantitative trait variation and expression variation are related. scRNAseq represents an appealing approach to eQTL mapping as it is possible to simultaneously genotype individual cells and measure expression in the same cell. As eQTL mapping requires large sample sizes to identify statistical relationships, the use of scRNAseq is likely to dramatically increase the statistical power of such studies. However, there are several technical challenges associated with scRNAseq and the authors' study is focused on addressing those challenges. Most of the points raised by my review of the initial version have been addressed. However, one point remains and one additional point should be considered.

(1) Given that the authors overcame many technical and analytical challenges in the course of this research, the study would be greatly strengthened through analysis of at least one, and ideally several, more conditions which would expand the conclusions that could be drawn from the study and demonstrate the power of using scRNAseq to efficiently quantify expression in different environments.

(2) In this version the authors have introduced the use of data imputation using a published algorithm, DISCERN. This has greatly increased the variation explained by their model as presented in figure 3. However, it is possible that the explained variance is now an overestimation as a result of using the imputed expression data. I think that it would be appropriate to present figure 3 using the sparse data presented in the initial version of the paper and the newly presented imputed data so that the reader can draw their own conclusions about the interpretation.

<https://doi.org/10.7554/eLife.93906.2.sa2>

Reviewer #2 (Public review):

The authors now say the main take-home for their work is (1) they have established methods for linkage mapping with scRNA-seq and that these (2) "can help gain insights about the genotype-phenotype map at a broader scale." My opinion in this revision is much the same as it was in the first round: I agree that they have met the first goal, and the second theme has been so well explored by other literature that I'm not convinced the authors' results meet the bar for novelty and impact. To my mind, success for this manuscript would be to support the claim that the scRNA-seq approach helps "reveal hidden components of the yeast genotype-to-phenotype map." I'm not sure the authors have achieved this. I agree that the new Figure 3 is a nice addition-a result that apparently hasn't been reported elsewhere (30% of growth trait variation can't be explained by expression). The caveats are that this is a negative result that needs to be interpreted with caution; and that it would be useful for the authors to clarify whether the ability to do this calculation is a product of the scRNA-seq method per se or whether they could have used any bulk eQTL study for it. Beside this, I regret to say that I still find that the results in the revision recapitulate what the bulk eQTL literature has already found, especially for the authors' focal yeast cross: heritability, expression hotspots, the role of cis and trans-acting variation, etc.

Likewise, when in the first round of review I recommended that the authors repeat their analyses on previous bulk RNA-seq data from Albert et al., my point was to lead the authors

to a means to provide rigorous, compelling justification for the scRNA-seq approach. The response to reviewers and the text (starting on line 413) says the comparison in its current form doesn't serve this purpose because Albert et al. studied fewer segregants. Wouldn't down-sampling the current data set allow a fair comparison? Again, to my mind what the current manuscript needs is concrete evidence that the scRNA-seq method per se affords truly better insights relative to what has come before.

I also recommend that the authors take care to improve the main text for readability and professionalism. It would benefit from further structural revision throughout (especially in the figure captions) to allow high-impact conclusions to be highlighted and low-impact material to be eliminated. Figure 4 and the results text sections from line 319 onward could be edited for concision or perhaps moved to supplementary if they obscure the authors' case for the scRNA-seq approach. The text could also benefit from copy editing (e.g. three clauses starting with "while" in the paragraph starting on line 456; "od ratio" on line 415). I appreciate the authors' work on the discussion, including posing big picture questions for the field (lines 426-429), but I don't see how they have anything to do with the current scRNA-seq method.

<https://doi.org/10.7554/eLife.93906.2.sa1>

Author response:

The following is the authors' response to the original reviews.

Reviewer #1:

Minor

(MN1) The segregants should be referred to as F2 segregants as they are derived from an F1 cross.

We thank the reviewer for pointing out this important oversight. We indeed analyzed segregants of an F1 cross and have corrected this in the text.

(MN2) The connections to eQTLs in other organisms should be addressed in the introduction and conclusion. For example, in humans, there has been little evidence for trans eQTLs in contrast to what has been found in yeast.

We thank the reviewer for pointing this out and improved our introduction and conclusion with such connections.

(M3) The authors state that an advantage of scRNAseq over bulk is that it captures rare cell populations (line 79), but this advantage is not exploited in this study.

While we did not explicitly demonstrate the effect of using scRNA-seq on capturing variation in rare cell populations, the referenced literature (21, 40) provides evidence that pooled scRNA-seq captures important expression heterogeneity (which implicitly contains potentially rare expression states). In our study, this is leveraged on F2 segregants to assess expression variation within the same lineage (genotype). This impacts the partitioning of expression variance from genotype.

Thus, we mentioned this point to further support the choice of using scRNA-seq for this analysis and showed that even a few single cells enable the reconstruction of the genome and expression profile of rare cell types.

(MN4) The authors use ~5% of the lineages from the original study. There is no rationale for why this is an appropriate sample size. Is there an argument for using more cells in eQTL mapping or conversely could the authors ask if fewer cells would provide similar conclusions by downsampling?

Although scRNA-seq is highly scalable, it has limitations in terms of throughput. Indeed, a single library with 10x Genomics generates data in the order of 10^4 wellcovered cells. With these limitations, our choice of ~5% of the lineages of the original study stems from the need to recover the same lineage multiple times within these 10^4 cells (in our study, each lineage is recovered on average 4 times).

While it is possible to run multiple libraries and sequencing lanes, budget limitations prevent us from running more libraries, especially since we expect power to scale with the square-root of the number of lineages (there is diminishing returns).

(MN5) I do not agree that the use of UMIs overcomes the challenges of low sequencing depth. UMIs mitigate the possible technical artifacts due to massive PCR amplification.

We thank the reviewer for this comment and will clarify this in the manuscript. Indeed, we intended to refer to the breadth of coverage (instead of the depth), which would usually manifest with massive PCR amplification of few transcripts.

(MN6) There is an inadequate reference to prior work on scRNAseq in yeast that established the methods used by the authors and eQTL mapping in human cells using scRNAseq.

We thank the reviewer for this and have added more context on scRNA-seq methods benchmark in yeast (drop-seq etc) and sc-eQTL in human. Additionally, we have cited Jariani et al. (2020) in eLife where similar techniques were employed for scRNA-seq in yeast.

(MN7) The use of empty quotes in Figure 4A is confusing and an alternative presentation method should be used.

We will remove these empty quotes characters and replace them with a more meaningful representation like “none”.

(MN8) The authors speculate about the use of predicted fitness instead of observed fitness, but this is something they could explicitly address in their current study.

We thank the reviewer for this comment but have decided not to perform a whole new bulk-segregant analysis experiment (X-QTL) to identify QTL that way. However, we do agree that we could in principle use the QTL that were identified in our previous study (Nguyen Ba et al, 2022). Despite this, we do not see the need for this because the predicted fitness is the overlap between genotype and phenotype (within the variance partitioning framework, it is the ‘narrow-sense heritability’ if one ignores epistasis). Thus, the use of predicted fitness when partitioning for expression variation would be constrained to that overlap (as opposed to the real observed fitness). This means that within the variance partitioning framework, the overlap of genotype, expression, and fitness is fully recapitulated by using predicted fitness instead (given that this predicted fitness is accurate to the narrow-sense heritability). In our previous study, we found that the QTL essentially predict all of the narrow-sense heritability. We believe it is therefore evident that the use of predicted fitness would be sufficient if and only if the expression variation independent of genotype is not associated with observed fitness.

We note that our study cannot generalize whether the overlap between genotype and expression fully captures fitness variation explained by expression. Indeed, we believe this is not generalizable to many other contexts (for example, in development). Thus, at present, the use of predicted fitness remains a speculation.

Major:

(MJ1) There is insufficient information provided about the nature of data. At a minimum, the following information should be provided to enable assessment of the study: What is the total library size, how many genes are identified per cell, how many UMIs are found per cell, what is the doublet rate, and how are doublets identified (e.g. on the basis of heterozygous calls at polymorphic loci?), how many times is each genotype observed, and how many polymorphic sites are identified per cell that are the basis of genotype inferences?

We understand that these metrics are relevant to the reader to have an idea of the power of our approach and integrate them in the manuscript in Table 1.

(MJ2) The prior study analyzed 18 different conditions, whereas this study only assays expression in a single condition. However, the power of the authors' approach is that its efficiency enables testing eQTLs in multiple conditions. The study would be greatly strengthened through analysis of at least one more condition, and ideally several more conditions. The previous fitness study would be a useful guide for choosing additional conditions as identifying those conditions that result in the greatest contrasts in fitness QTL would be best suited to testing the generalizations that can be drawn from the study.

We agree that a major strength of our approach is that it rapidly allows eQTL mapping in several conditions. While the experiments presented here are likely less expensive than the classical eQTL mapping experiments, the cost of 10x genomics and sequencing is still an important consideration. The pleiotropy analysis of the prior study was substantially difficult to interpret and put in context, and thus we decided to focus on a proof of concept and leave room for a more thorough analysis of multiple environments for a future study. We acknowledge that this is a main weakness of our manuscript.

(MJ3) Alternatively, the authors could demonstrate the power of their approach by applying it to a cross between two other yeast strains. As the cross between BY and RM has been exhaustively studied, applying this approach to a different cross would increase the likelihood of making novel biological discoveries.

We thank the reviewers for this suggestion, and it is indeed something that our lab is considering. Currently, one of our main point of the manuscript still relies on growth measurements of segregants (the fitness), which we cannot obtain from segregants and scRNA-seq alone.

Unfortunately, in this experimental design, it is difficult to obtain the fitness of cells and the genotype simultaneously because the barcode of the segregant is not expressed and not frequently read during genotyping. Thus, we still need to perform a whole QTL panel for a new cross without substantial re-engineering.

That being said, we are working on this but feel that including a new panel in this study is beyond the scope of our manuscript.

(MJ4) Figure 1 is misleading as A presents the original study from 2022 without important details such as how genotypes were identified. It is unclear what the barcode is in this study and how it is used in the analysis. Is the barcode for each lineage transcribed so that it is identified in the scRNA-seq data? Or, does the barcode in B refer to the cell index barcode? A clearer presentation and explanation of terms are needed to understand the method.

Because F2 segregant lineage barcodes are not expressed, the barcode indicated in Figure 1B refers to cell barcodes from 10x Genomics. Our present study does not make use of the lineage barcode. We clarified this in the figure clarifying that panel A refers to the original study from 2022 and explicitly mentioning ‘cell barcodes’.

(MJ5) The rationale for the analysis reported in Figure 2B is unclear. The fitness data are from the previous study and the goal is to estimate the heritability using the genotyping data from the scRNA-Seq data. What is the explanation for why the data don't agree for only one condition, i.e. 37C? And, what are we to understand from the overall result?

The rationale of Figure 2A/B is to show that cell lineage genotyping with scRNA-seq yields consistent results with previous genotype-phenotype analyses of the same cross. While Figure 2A shows that the single-cell imputed genotypes resemble the reference panel (sequenced in the Nguyen Ba 2022 study), Figure 2B shows that the variance partitioning to associate genotype to phenotype can be performed using the single-cell genotypes themselves (bypassing the reference panel). We believe this is an interesting result given that the reads obtained by scRNA-seq are constrained to a subset of SNP. However, we note that if the imputed single-cell genotypes were perfectly matching with the reference panel, it would not be surprising that one could do genotype-phenotype mapping from the single-cell genotypes.

In Figure 2B, we tested whether the similarity of the single-cell imputed genotypes to the reference panel was enough to estimate heritabilities (another summary statistic).

In the remaining paragraphs of that result section, we further discuss that the single-cell lineage genotypes can be used for QTL mapping as well, recapitulating many of the QTL identified in the reference panel (provided that one controls for power). This result did not make it as a main Figure but is included in Figure S4.

That being said, we decided to update the figure by comparing the estimates in subsamples of batch1 scRNA-seq to subsamples of batch 1 reference panel and subsamples of the full reference panel. Subsamples were performed to control for power in the variance partitioning. We also noticed that the fitness of several F2 segregants is missing for the phenotypes 33C, 35C and 37C in the original study so we decided to exclude these environments.

(MJ6) Figure 3 presents an analysis of variance partitioning as a Venn diagram. This summarized result is very hard to understand in the absence of any examples of what the underlying raw data look like. For example, what does trait variation look like if only genotype explains the variance or if only gene expression explains the variance? The presented highly summarized data is not intuitive and its presentation is poor - the result that is currently provided would be easier to read in a table format, but the reader needs more information to be able to interpret and understand the result.

The Venn diagram is largely adopted in the context of variance partitioning (see Cohen, Jacob, and Patricia Cohen. 1975. *Applied Multiple Regression/Correlation Analysis for the Behavioral Sciences*.) but we realize that it has not been used often for displaying heritability estimates.

To this end, we have added explanatory labels for the biological meaning of the areas or components of the diagram in the Figure and in the text.

(MJ7) I am concerned about the conclusions that can be drawn about expression heritability. The authors claim that expression heritability is correlated with expression levels. It seems likely that this reflects differing statistical power. How can this possibility be excluded?

We thank the reviewer for highlighting this. We now explicitly acknowledge this potential confounding factor in the manuscript.

(MJ8) Conversely, the authors claim that the genes with the lowest heritability are genes involved in the cell cycle. However, uniquely in scRNA-seq, cell cycle regulated genes appear to have the highest variance in the data as they are only expressed in a subset of cells. Without incorporating this fact one would erroneously conclude that the variation is not heritable. To test the heritability of cell cycle regulation genes the authors should partition the cells into each cell cycle stage based on expression.

The reviewer is right to say that the low heritability of cell cycle control genes could be explained by the fact that these genes are only expressed in a subset of the dataset. Indeed, a high transcriptomic variance does not necessarily imply a low expression heritability: the cell cycle could be the residual of the expression heritability model, i.e. it explains expression variance with low association to genetic mutation.

That being said, our result is consistent with results obtained from yeast bulk RNA-seq (Albert et al. 2018), in which cell cycle is averaged out.

In our study, we also average out the cell-cycle as we use the consensus expression and the consensus genome to estimate the heritability.

(MJ9) I do not understand Figure S5 and how eQTL sites are assigned to these specific classes given that the authors say that causative variation cannot be resolved because of linkage disequilibrium.

The rationale for Figure S5 is to show that the QTL model obtained from single-cell data is consistent with the reference panel QTL mapping experiment. Although there is uncertainty around the exact position of the QTL, we relied on the loci with the highest likelihood and showed that the datasets have consistent features. This is enabled by the fact that the QTL identified using the scRNA-seq genotypes are the ones with largest effect size in the reference panel, and are thus more likely to be mapped accurately.

(MJ10) The paragraph starting at line 305 is very confusing. In particular, the authors state that they identify a hotspot of regulation at the mating type locus. It is not obvious why this would be the case. Moreover, they claim that they find evidence for both MATa and MATalpha gene expression. Information is not provided about how segregants were isolated, but assuming that the authors did not dissect 25,000 tetrads to obtain 100,000 segregants I would infer that random spore using SGA was used. In that case, all cells should be MATa. The authors should clarify and explain this observation.

Although most of the cells have the MATa mating type (as selected by random spore using SGA), it is well known and discussed in Nguyen Ba et al. paper that there are few lineages with other mating types or diploids (they are leakers in the selection process).

Indeed, we verified that we can detect a small number of MATalpha cells or diploids within this pool.

(MJ11) Ultimately, it is not clear what new biological findings the authors have made. There are no novel findings with respect to causative variation underlying eQTLs and I would encourage the authors to make clearer statements in their abstract, introduction, and conclusion about the key discoveries. E.g. What are the "new associations between phenotypic and transcriptomic variations" mentioned in the abstract?

This paper focuses more on the proof of concept that scRNA-seq can help integrate expression data in GPM analysis to reveal broad scale associations between fitness and expression. Indeed, novel findings include new hotspots of expression regulation in the RM/BY genetic background, we find that trans-regulation of expression has more impact than cis-regulation on fitness and evaluate the strength of the association between the genome, the transcriptome and fitness (in one environment). Additionally, the analysis reveals biological questions that cannot be answered even by increasing the experimental scale of eQTL mapping experiments. For example, we find that most of the missing heritability is not explained by expression. These key points will be clarified in the abstract, introduction and conclusion as suggested by the editors.

Reviewer #2:

(MJ1) Most of the figures center on methods development and validation for the authors' single-cell RNA-seq in the yeast cross [...] One potential novelty of the study is the methods per se: that is, showing that scRNA-seq works for concomitant genotyping and gene expression profiling in the natural variation context. The authors' rigor and effort notwithstanding: in my view, this can be described as modest in terms of principles. That is, the authors did a good job putting the scRNA-seq idea into practice, but their success is perhaps not surprising or highly relevant for work outside of yeast (as the discussion says).

Although the scope of the method is limited, we think that it can apply to any largescale dataset in which transcription variance and genetic diversity are not small. This can help reduce the lack of associations between trait heritability and expression regulation, which is frequent as these two parameters are often not measured within the same dataset.

We can, however, think of some other settings where a similar experiment may be interesting. This includes, for example, pooling cells from different human individuals (with enough genetic diversity) and applying the same scRNA-seq method to back-identify the individuals and matching them to a particular phenotype. We believe our proof of concept is therefore an important contribution as these other experiments might have broad implications.

(MJ2) The more substantive claim by the authors for the impact of the study is that they make new observations about the role of expression in phenotype (lines 333-335). The major display item of the manuscript on this theme is Figure 4A, reporting which loci that control growth phenotype (from an earlier paper) also control expression. This is solid but I regret to say that the results strike me as modest.

This paper focuses more on the proof of concept that scRNA-seq can help integrate expression data in GPM analysis to reveal broad scale associations between fitness and expression. Indeed, novel findings include new hotspots of expression regulation in the RM/BY genetic background, we find that trans-regulation of expression has more impact than cis-regulation on fitness and evaluate the strength of the association between the genome, the transcriptome and fitness (in one environment). Additionally, the analysis reveals biological questions that cannot be answered even by increasing the experimental scale of eQTL mapping experiments. For example, we find that most of the missing heritability is not

explained by expression. These key points will be clarified in the abstract, introduction and conclusion as suggested by the editors.

(MJ3) The discussion makes some perhaps fairly big claims that the work has helped "bridge understanding of how genetic variation influences transcriptomic variation" and ultimately cellular phenotype. But with the data as they stand, the authors have missed an opportunity to crystallize exactly how a given variant affects expression (perhaps in waves of regulators affecting targets that affect more regulators) and then phenotype, except for the speculations in the text on lines 305-319. The field started down this road years ago with Bayesian causality inference methods applied to eQTL and phenotype mapping (via e.g. the work of Eric Schadt). The authors could now try Mendelian randomization-type fine-grained detailed models for more firepower toward the same end, and/or experimental tests of the genotype-to-expression-to-phenotype relationship. I would see these directions, motivated by fundamental questions that are relevant to the field at large, as leading to a major advance for this very crowded field. As it stands, I felt their absence in this manuscript especially if the authors are selling principles about linking expression and phenotype as their take-home.

We thank the reviewer for this suggestion and agree that the analysis of the genotypeto-expression-to-phenotype relationship would benefit from a more fine-grain model. While we are interested in exploring this, we decided to limit the scope of this manuscript to the proof of concept that scRNA-seq can help gain insights about the genotypephenotype map at a broader scale.

(MN1) I also wonder whether the co-mapping of expression and growth traits in Figure 4A would have been possible with e.g. the bulk RNA-seq from Albert et al., 2018, and I recommend that the authors repeat the Figure 4A-type analyses with the latter to justify their statement that their massive scRNA data set would actually be necessary for them to bear fruit (lines 386-388).

By repeating our eQTL hotspot analysis with Albert et al. (2018) data, we observed a non-significant association between eQTL hotspot and QTL (χ^2 $p = 0.50$). That being said, there are some differences in the Albert et al. Experiment that preclude us from conclusively saying whether the bulk RNA-seq experiments by Alberts would not bear fruit. Indeed, that experiment is only 4 times smaller in scale and so we would not expect dramatic differences. To highlight power differences, the Albert et al. Paper identified about 6 eQTL per gene, while our study identified about 21 which is consistent with the power differences.

This highlights that this scRNA-seq experiment is scalable, so the technique may be useful for further studies. In addition, this pooled scRNA-seq strategy enables analysis of the association of transcription with phenotype.

(MN2) I also read the discussion of the manuscript as bringing to the fore some of the challenges a reader has in judging the current state of the results to be of actionable impact. The discussion, and the manuscript, will be improved if the authors can put the work in context, posing concrete questions from the field and stating how they are addressed here and what's left to do.

We agree with the reviewer and have summarized our answers to some of the questions in the field in the discussion section.

All that being said, we acknowledge the limitations of our study.

<https://doi.org/10.7554/eLife.93906.2.sa0>