

# **Compressed Perturb-seq: highly efficient screens for regulatory circuits using random composite perturbations**

**Authors:** Douglas Yao<sup>1</sup>, Loic Binan<sup>2</sup>, Jon Bezney<sup>2,13</sup>, Brooke Simonton<sup>2</sup>, Jahanara Freedman<sup>2</sup>, Chris J. Frangieh<sup>2,3</sup>, Kushal Dey<sup>4</sup>, Kathryn Geiger-Schuller<sup>5</sup>, Basak Eraslan<sup>5</sup>, Alexander Gusev<sup>2,6,7,15</sup>, Aviv Regev<sup>2,14,15</sup>, Brian Cleary<sup>8,9,10,11,12,15</sup>

## **Affiliations:**

1. Program in Systems, Synthetic, and Quantitative Biology, Harvard University, Cambridge, MA
2. Klarman Cell Observatory, Broad Institute of Harvard and MIT, Cambridge, MA
3. Department of Electrical Engineering and Computer Science, Massachusetts Institute of Technology, Cambridge, MA
4. Harvard T.H. Chan School of Public Health, Boston, MA
5. Genentech, South San Francisco, CA
6. Department of Medical Oncology, Dana-Farber Cancer Institute, Boston, MA
7. Division of Genetics, Brigham and Women's Hospital, Boston, MA
8. Faculty of Computing and Data Sciences, Boston University, Boston, MA
9. Department of Biology, Boston University, Boston, MA
10. Department of Biomedical Engineering, Boston University, Boston, MA
11. Program in Bioinformatics, Boston University, Boston, MA
12. Biological Design Center, Boston University, Boston, MA
13. Current address: Department of Genetics, Stanford University School of Medicine, Stanford, CA
14. Current address: Genentech, South San Francisco, CA
15. These authors jointly supervised this work

Lead contact: [bcleary@bu.edu](mailto:bcleary@bu.edu) (B.C.)

# **Abstract**

Pooled CRISPR screens with single-cell RNA-seq readout (Perturb-seq) have emerged as a key technique in functional genomics, but are limited in scale by cost and combinatorial complexity. Here, we reimagine Perturb-seq's design through the lens of algorithms applied to random, low-dimensional observations. We present compressed Perturb-seq, which measures multiple random perturbations per cell or multiple cells per droplet and computationally decompresses these measurements by leveraging the sparse structure of regulatory circuits. Applied to 598 genes in the immune response to bacterial lipopolysaccharide, compressed Perturb-seq achieves the same accuracy as conventional Perturb-seq at 4 to 20-fold reduced cost, with greater power to learn genetic interactions. We identify known and novel regulators of immune responses and uncover evolutionarily constrained genes with downstream targets enriched for immune disease heritability, including many missed by existing GWAS or trans-eQTL studies. Our framework enables new scales of interrogation for a foundational method in functional genomics.

# INTRODUCTION

Pooled perturbation screens with high-content readouts, ranging from single-cell RNA-seq (Perturb-seq)<sup>1-4</sup> to imaging-based spatial profiling<sup>5-7</sup>, are now enabling systematic studies of the regulatory circuits that underlie diverse cell phenotypes. Perturb-seq has been applied to various model systems, leading to insights about diverse cellular processes including the innate immune response<sup>2</sup>, *in vivo* effects of autism risk genes in mice<sup>8</sup> and organoids<sup>9,10</sup>, and genome-scale effects on aneuploidy, differentiation, and RNA splicing<sup>11</sup>. Integrating data from population-level genetic screens has also elucidated human disease mechanisms<sup>12</sup>.

However, due to the large number of genes in the genome, large-scale Perturb-seq screens are still prohibitively expensive and are often limited by the number of available cells, especially for primary cell systems<sup>13</sup> and *in vivo* niches<sup>8</sup>. In addition, the exponentially larger number of possible genetic interactions makes it impossible to conduct exhaustive combinatorial screens for genetic interactions using existing approaches, so current Perturb-seq studies of genetic interactions are very modest and focused<sup>14</sup>. Several approaches have been developed to improve the efficiency of scRNA-seq and/or Perturb-seq, including overloading droplets with multiple pre-indexed cells (SciFi-seq<sup>15</sup>) or pooling multiple guides within cells<sup>16</sup>. However, pre-indexing requires an additional laborious and complex experimental step, while guide pooling has only been used to study *cis* and not *trans* effects of perturbations.

We propose an alternative approach to greatly increase the efficiency and power of Perturb-seq for both single and combinatorial perturbation screens, inspired by theoretical results from compressed sensing<sup>17-19</sup> that apply to the sparse and modular nature of regulatory circuits in cells. To elaborate, perturbation effects tend to be *sparse* in that most perturbations affect only a small number of genes or co-regulated gene programs<sup>2</sup>. In this scenario, rather than assaying each perturbation individually, we can measure a much smaller number of *random combinations* of perturbations (forming what we call “composite samples”) and accurately learn the effects of individual perturbations from the composite samples using sparsity-promoting algorithms. Moreover, with certain types of composite samples, we can efficiently learn both first-order effects (*i.e.*, from single gene perturbations) and higher-order genetic interaction effects from the same data. We have previously shown that experiments that measure random compositions of the underlying biological dataset can greatly increase the efficiency of measuring expression profiles<sup>20</sup> and imaging transcriptomics<sup>21</sup>.

Here, we develop two experimental strategies to generate composite samples for Perturb-seq screens, and we introduce an inference method, Factorize-Recover for Perturb-seq analysis (FR-Perturb), to learn individual perturbation effects from composite samples. We apply our approach to 598 genes in a human macrophage cell line treated with lipopolysaccharide (LPS). By comparing compressed Perturb-seq to conventional Perturb-seq conducted in the same system, we demonstrate the enhanced efficiency and power of our approach for learning single perturbation effects and second-order genetic interactions. We derive insights into immune regulatory functions and illustrate their connection to human disease mechanisms by integrating data from genome-wide association studies (GWAS) and expression quantitative trait loci (eQTL) studies.

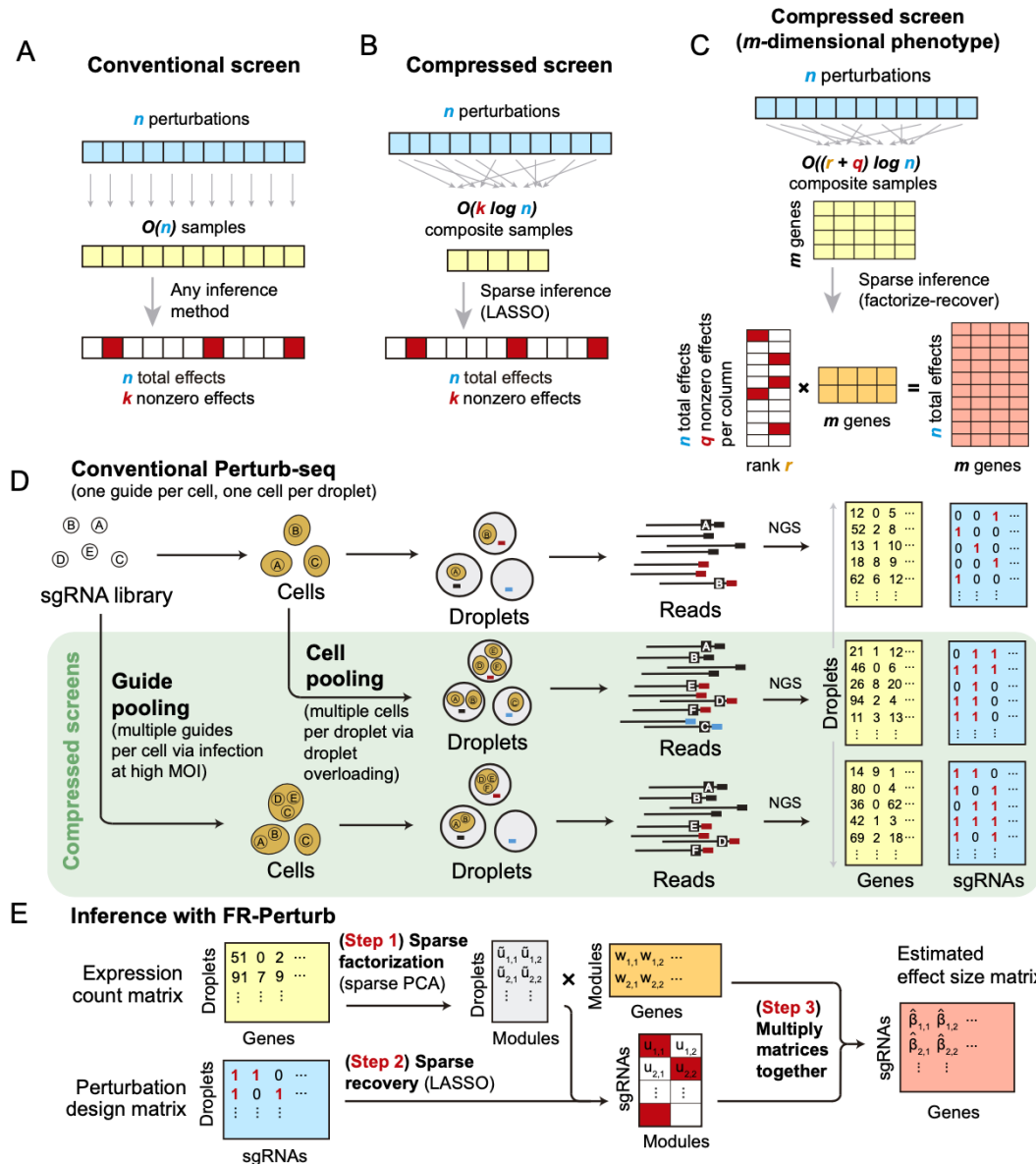
## RESULTS

### A compressed sensing framework for perturbation screens

In conventional Perturb-seq, each cell in a pool receives one or more genetic perturbations. Each cell is then profiled for the identity of the perturbation(s) and the expression levels of  $m \approx 20,000$  expressed genes. Our goal is to infer the effect sizes of  $n$  perturbations on the phenotype, which can be the entire gene expression profile ( $n \times m$  matrix) or an aggregate multi-gene phenotype<sup>2,3,11</sup> such as an expression program or cell state score (length- $n$  vector). In both cases, we need  $O(n)$  samples to learn the effects of  $n$  perturbations (**Figure 1A**) (where sample replicates introduce a constant factor that is subsumed under the big O notation), such that the number of samples scales *linearly* with the number of perturbations.

Based on the theory of compressed sensing<sup>17</sup>, there exist conditions under which far fewer than  $O(n)$  samples are sufficient to learn the effects of  $n$  perturbations. In general, if the perturbation effects are sparse (*i.e.*, relatively few perturbations affect the phenotype), or are sparse in a latent representation (*i.e.*, perturbations tend to affect relatively few latent factors that can be combined to “explain” the phenotype), then we can measure a small number of random composite samples (comprising *linear combinations* of individual sample phenotypes) and decompress those measurements to infer the effects of individual perturbations. Composite samples can be generated either by randomly pooling perturbations in individual cells, or by randomly pooling cells containing one perturbation each (see below).

The number of required composite samples depends on whether the phenotype is single-valued or high-dimensional. When the phenotype is single-valued (*e.g.*, fitness),  $O(k \log n)$  composite samples suffice to accurately recover the effects of  $n$  perturbations<sup>18,19</sup>, where  $k$  is the number of nonzero elements among the  $n$  perturbation effects (**Figure 1B**). When most genes do not affect the phenotype,  $k$  grows more slowly than  $n$ , and the number of required composite samples scales logarithmically or at worst sub-linearly with the number of perturbations. Meanwhile, when the phenotype is an  $m$ -dimensional gene expression profile, an efficient approach involves inferring effects on latent expression factors, then reconstructing the effects on individual genes from these factors using the “factorize-recover” algorithm<sup>22</sup>. This approach requires  $O((q + r) \log n)$  composite samples, where  $r$  is the *rank* of the  $n \times m$  perturbation effect size matrix (*i.e.*, the maximum number of its linearly independent column vectors), and  $q$  is the maximum number of nonzero elements in any column of the left matrix of the factorized effect size matrix (**Figure 1C**). In our case,  $r$  is the number of distinct groups of “co-regulated” genes whose expression changes concordantly in response to any perturbation, while  $q$  is the maximum number of “co-functional” perturbations with nonzero effects on any individual module. Due to the modular nature of gene regulation<sup>23,24</sup>,  $r$  and  $q$  are expected to remain small when  $n$  increases. Indeed, we observed a relatively small number of co-functional and co-regulated gene groups (small  $q$  and  $r$ , respectively, relative to  $n$ ) in previous Perturb-seq screens in various systems<sup>2,13</sup>. Thus, the number of required composite samples will scale logarithmically or at worst sub-linearly with  $n$ , leading to much fewer required samples than the conventional approach with large  $n$ .



**Figure 1. Framework for compressed Perturb-seq.** (A) Schematic for conventional perturbation screen with single-valued phenotype. Each sample (yellow) receives a single perturbation (blue). The required number of samples scales linearly with the number of perturbations, as captured by the  $O(n)$  term. (B) Schematic for compressed perturbation screen with single-valued phenotype. Each “composite” sample (yellow) represents a random combination of perturbations (blue). The required number of samples scales sub-linearly with the number of perturbations given the following: (1) the effects of the perturbations are sparse (*i.e.*,  $k$  increases more slowly than  $n$ ), and (2) sparse inference (typically LASSO) is used to infer the effects from the composite sample phenotypes. (C) Schematic for compressed perturbation screen with high-dimensional phenotype, which is the main use case for Perturb-seq. The required number of samples scales sub-linearly with the number of perturbations given the following: (1) the effects of the perturbations are sparse and act on a relatively small number of groups of correlated genes (*i.e.*,  $q$  and  $r$  increase more slowly than  $n$ ), and (2) sparse inference (namely the “factorize-recover” algorithm<sup>22</sup>) is used to infer the effects from the composite

sample phenotypes. **(D)** Two experimental strategies for generating composite samples for Perturb-seq. Both “cell pooling” and “guide pooling” change one step of the conventional Perturb-seq protocol. The result is a sample whose phenotype corresponds to a random linear combination of the phenotypes of samples from the conventional Perturb-seq screen. **(E)** Schematic of computational method used to infer perturbation effects from composite sample phenotypes, based on the “factorize-recover” algorithm<sup>22</sup>.

### Experimentally generating composite samples for compressed Perturb-seq

We generated composite samples for compressed Perturb-seq by either randomly pooling cells containing one perturbation each in overloaded scRNA-seq droplets<sup>15</sup>, or by randomly pooling guides in individual cells via infection with a high multiplicity of infection (MOI)<sup>2,16</sup> (**Figure 1D**). Either method can be used to learn the same underlying perturbation effects, but each has different strengths and limitations (**Discussion**).

For pooling cells with droplet overloading, we load cells into the microfluidics chip at much higher concentration than usual so that each droplet contains multiple cells. The resulting expression counts from a given droplet are proportional to the average expression counts of the cells in the droplet, while the set of detected guides is the union of those present in any of the constituent cells. We do not intend to match reads to specific cells within each droplet and in fact cannot do so without additional indexing<sup>15</sup>. Rather, we model the expression counts of each gene as the expected counts from an unperturbed (control) cell multiplied by the average of the fold-change effects from each guide in the droplet (**Methods**).

For pooling guides, we transduce cells at high MOI so that multiple random guides enter each cell. We then measure the expression counts from individual cells and identify their constituent guides. We model the log-expression counts of each gene as the log of the expected counts from a control cell plus the sum of log fold-change effects from the guides in the cell (**Methods**). This model, where observations are a random linear combination of log fold-change effect sizes from different guides, makes the non-trivial assumption that the effect sizes of guides tend to combine additively in log expression space when present in the same cell. We can test the validity of this assumption by looking at the concordance between effect size estimates from cells with only one guide versus cells with multiple guides (see below). Pooling guides has a key benefit over pooling cells, in that the generated data can be used to estimate both first-order effects and higher-order genetic interactions (with appropriate sample sizes and explicit interaction terms in the model) (**Methods**). We illustrate the feasibility of estimating second-order effects from our guide-pooled data below.

### FR-Perturb infers effects from compressed Perturb-seq

To infer perturbation effects from the composite samples, we devised a method called FR-Perturb based on the “factorize-recover” algorithm<sup>22</sup> (**Methods**). FR-Perturb first factorizes the expression count matrix with sparse factorization (*i.e.*, sparse PCA), followed by sparse recovery (*i.e.*, LASSO) on the resulting left factor matrix comprising perturbation effects on the latent factors. Finally, it computes perturbation effects on individual genes as the product of the left factor matrix from the recovery step with the right factor matrix (comprising gene weights in each latent factor) from the first factorization step (**Figure 1E**; **Methods**). We obtained p-values and false discovery rates (FDR) for all effects by permutation testing (**Methods**). We evaluated FR-Perturb by comparing

it to existing inference methods for Perturb-seq, namely elastic net regression<sup>2</sup> and negative binomial regression<sup>16</sup> (see below).

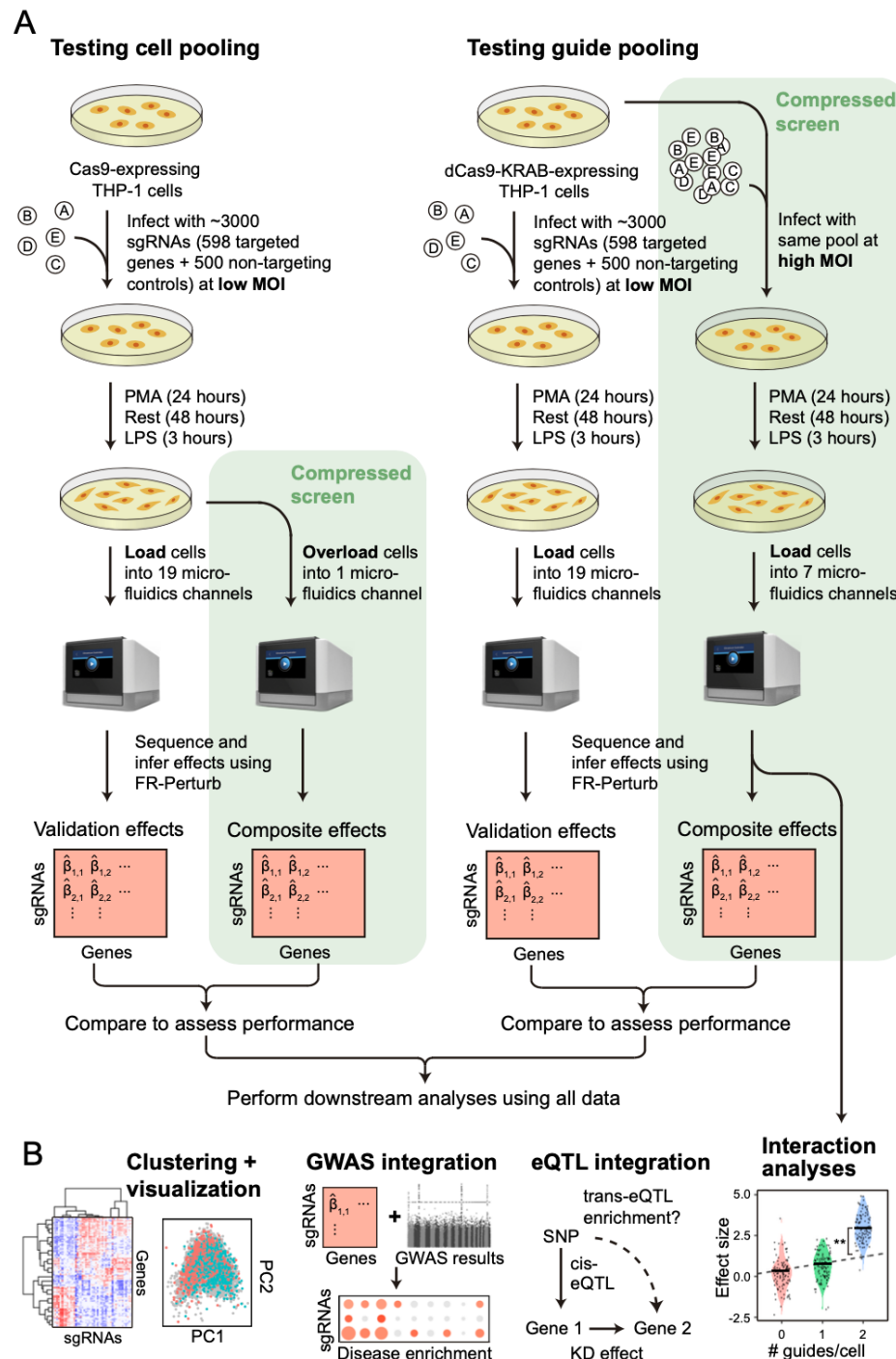
### Compressed Perturb-seq screens of the LPS response

We implemented and evaluated compressed Perturb-seq in the response of THP1 cells (a human monocytic leukemia cell line) to stimulation with LPS when either pooling cells or pooling guides (**Figure 2A-B**). In each case, we also performed conventional Perturb-seq, targeting the same genes in the same system for comparison. We selected 598 genes to be perturbed from seven mostly non-overlapping immune response studies (**Table S1**), including genes with roles in the canonical LPS response pathway (34 genes), GWAS for inflammatory bowel disease (79 genes) and infection (106 genes), Mendelian immune diseases from OMIM with keywords for “bacterial infection” (85 genes) and “NF-kappa-B” (102 genes), a previous genome-wide screen for effects on TNF expression in mouse BMDCs<sup>25</sup> (93 genes), and genes with large genetic effects in *trans* on gene expression from an eQTL study in patient-derived macrophages stimulated with LPS<sup>26</sup> (79 genes) (**Methods**, **Figure S1A**). We designed 4 sgRNAs for each gene and 500 each of non-targeting or safe-targeting control sgRNAs, resulting in a total pool of 3,392 sgRNAs (**Methods**). We introduced the sgRNAs into THP1 cells via a modified CROP-seq vector<sup>4</sup> (**Methods**). After transduction and selection, we treated cells with PMA for 24 hours and grew them for another 48 hours as they differentiated into a macrophage-like state<sup>27</sup>, then treated them with LPS for three hours before harvesting for scRNA-Seq (**Methods**). For our cell-pooled screen, we used CRISPR-Cas9 to knock out genes<sup>2</sup>, whereas for our guide-pooled screen we used CRISPRi with dCas9-KRAB to knock down gene expression<sup>1</sup> (**Figure 2A**) to avoid cellular toxicity due to multiple double-stranded breaks in individual cells<sup>28</sup>.

By design, the two compressed screens were substantially smaller than their corresponding conventional screens. In the cell-pooled screen, we analyzed a single channel of droplets (10x Genomics, **Methods**) overloaded with 250,000 cells, while for the corresponding conventional Perturb-seq screen we analyzed 19 channels at normal loading. We sequenced the library from the overloaded channel to a depth of 4-fold more reads than a conventional channel to account for the larger number of non-empty droplets and greater expected RNA content per droplet. After quality control, there were 32,700 droplets containing at least one sgRNA from the overloaded channel (vs. 4,576 droplets/channel for a total of 86,954 droplets from the conventional screen) (**Figure 3A**), with a mean of 1.86 sgRNAs per non-empty droplet (conventional: 1.11) (**Figure 3B**) and a mean of 90 droplets containing a guide for each perturbed gene (conventional: 144) (**Figure 3C**). We observed 14,987 total genes with measured expression (conventional: 17,552). Thus, the cell-pooled screen had >7 times the number of non-empty droplets per channel compared to the conventional screen; considering library preparation and sequencing costs, it was approximately 8 times cheaper.

In the guide-pooled experiment, we infected cells expressing dCas9-KRAB at high MOI (**Methods**) and profiled a single cell in each droplet across seven channels, while for the corresponding conventional Perturb-seq we infected cells with the same guide library at low MOI and analyzed 19 channels. From the guide-pooled experiment, we obtained 24,192 cells after filtering (conventional: 66,283), where 35% of the cells (8,448) contained three or more guides (**Figure 4A**), with 2.50 guides on average per cell (conventional: 1.13) (**Figure 4B**) and 101 cells containing a guide for each perturbed gene on average (conventional: 115) (**Figure 4C**). We

measured expression for 16,268 total genes (conventional: 18,617). The guide-pooled screen was approximately 3 times cheaper than the conventional screen.



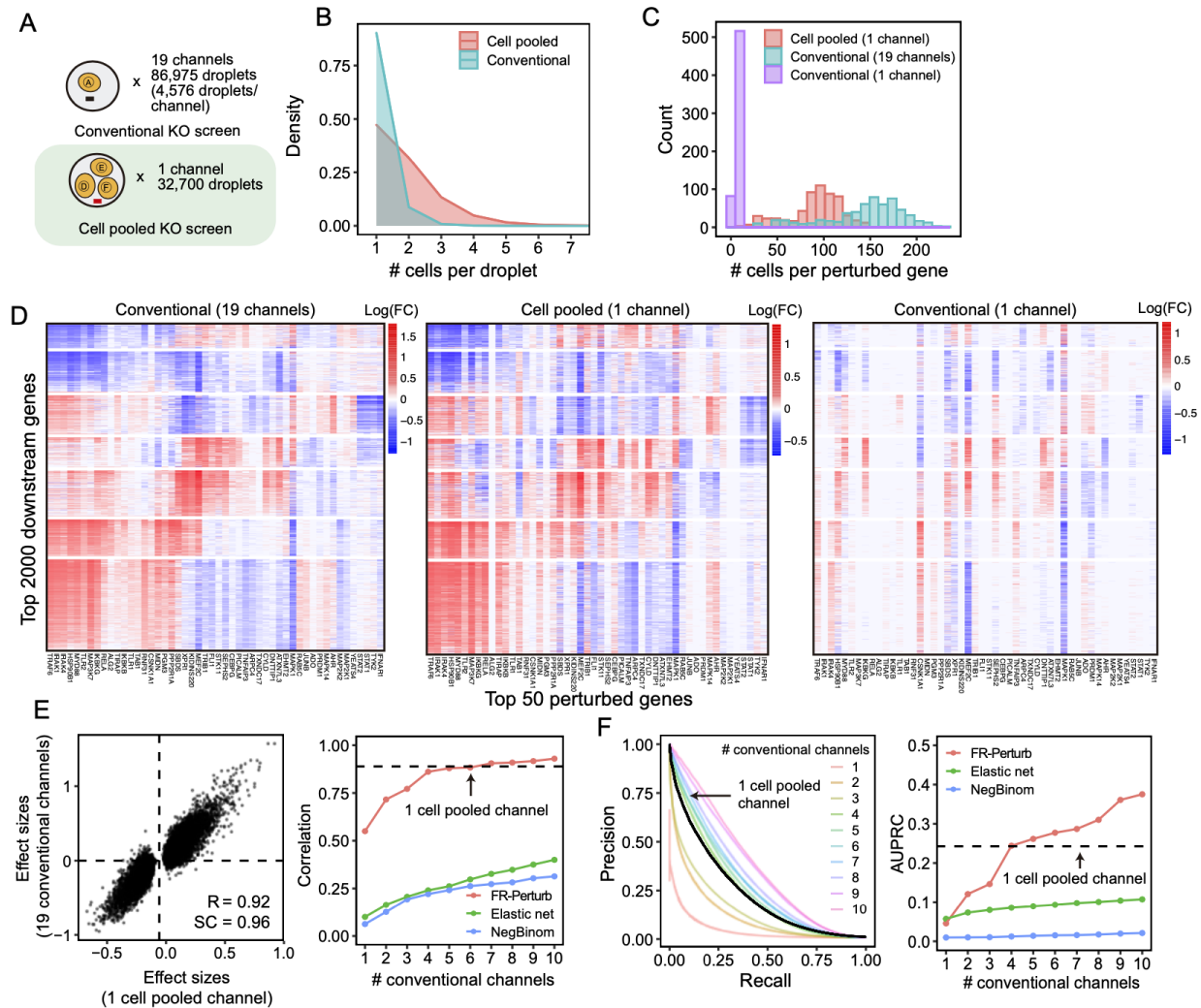
**Figure 2. Experimental overview.** (A) Outline of experiments used to test and validate cell pooling (left) and guide pooling (right). (B) Downstream analyses performed using perturbation effects from all experiments.

## Compressed Perturb-Seq by cell pooling yields accurate perturbation effects at high efficiency

The perturbation effect sizes estimated by Perturb-FR from the cell-pooled Perturb-seq screen (**Methods**) agreed well with its conventional counterpart (**Table S2**). When estimating effects, we included read count, cell cycle, and proportion of mitochondrial reads as covariates<sup>2</sup>, and we combined sgRNAs targeting the same gene while retaining the subset of sgRNAs for a gene with maximal concordance of effects across random subsets of the data (**Methods**). The significant effects from the compressed experiment (N = 19,909) were strongly correlated with the corresponding effects from the conventional experiment (Pearson's R = 0.92, sign concordance = 0.96, **Figure 3E**). Notably, we observed many more significant effects overall in the conventional screen than the cell-pooled screen (216,220 vs. 19,909; FDR q-value < 0.05), but this is expected given that we intentionally generated a much larger and more highly powered conventional screen to enable data splitting and cross validation analyses (see below).

The cell-pooled experiment yielded substantially more signal per experimental unit (channel) than the conventional one (**Figure 3D-F**). First, the global clustering of effects learned from a single cell-pooled channel was much less noisy than from a single conventional channel (**Figure 3D**). Moreover, approximately four conventional channels were needed to obtain the same number of significant effects as one cell-pooled channel (**Figure S2A**). Next, to quantitatively assess the specificity of each approach, we held out half of the conventional data as a validation set, then we down-sampled the remaining half to different numbers of channels and compared the top 19,909 most significant effects learned from the down-sampled data (matching the number of significant effects in the cell-pooled screen) to those in the held-out validation set. We found that 5-6 conventional channels were needed to achieve equivalent validation accuracy (correlation) as one cell-pooled channel (**Figure 3E**). The relative efficiency gains of the compressed screen were consistent when varying the number of effects being compared (**Figure S2C**). We also assessed the sensitivity of each approach by testing whether the significant effects determined from the validation set were recovered by the down-sampled conventional or cell-pooled screens. We constructed precision-recall curves, calling “true positives” the 79,100 significant effects from the validation dataset and varying the classification threshold by the significance of the effects in the down-sampled conventional or cell-pooled datasets. One cell-pooled channel had comparable AUPRC to 4 conventional channels (**Figure 3F**), with consistent efficiency gains when varying the number of true positive effects (**Figure S2C**).

Moreover, FR-Perturb substantially outperformed the established inference methods we tested: elastic net regression<sup>2</sup> and negative binomial regression<sup>16</sup>. Repeating the same analyses as above with each method (**Methods**), the concordance between the down-sampled conventional data and validation data, and between cell-pooled and conventional data, was much higher with FR-Perturb than prior methods (**Figure 3E-F**, **Figure S2D**). By down-sampling the cell-pooled screen, we found that ~1/5 of a cell-pooled channel analyzed with FR-Perturb achieved the same validation accuracy as 10 conventional channels analyzed with existing methods (**Figure S2B**). We assess the cost savings of cell pooling over the conventional approach while factoring in sequencing costs in the **Discussion**.



**Figure 3. Evaluating cell-pooled Perturb-seq versus conventional Perturb-seq.** (A) Number of channels and droplets from the conventional validation screen (top) and cell-pooled screen (bottom). (B) Distribution of droplets based on number of cells they contain for the cell-pooled and conventional screens. In practice, we only directly measure the number of guides/droplet rather than cells/droplets, but these quantities are equivalent given 1 guide/cell. (C) Distribution of the number of cells containing a guide targeting each perturbed gene in the cell-pooled screen and conventional screen. For the conventional screen, we show the distributions for both the full screen (19 channels) and an equivalent number of channels as the cell-pooled screen (1 channel). (D) Heatmaps of the top effect sizes (inferred with FR-Perturb) from the conventional screen (left), with the same effect sizes shown for the cell pooled screen (middle) and one equivalent channel of the conventional screen (right). X-axis: top 50 perturbed genes, based on their average magnitude of effect on all 17,552 downstream genes. Y-axis: top 2,000 downstream genes, based on the average magnitude of effects of all 598 perturbed genes acting on them. Rows and columns are clustered based on hierarchical clustering in the leftmost plot. For the left plot, all effects with FDR q-value > 0.2 are whited out, whereas this threshold is relaxed to 0.5 for the middle and right plots. (E) (Left) Scatterplot of all significant effects (FDR q < 0.05; N =

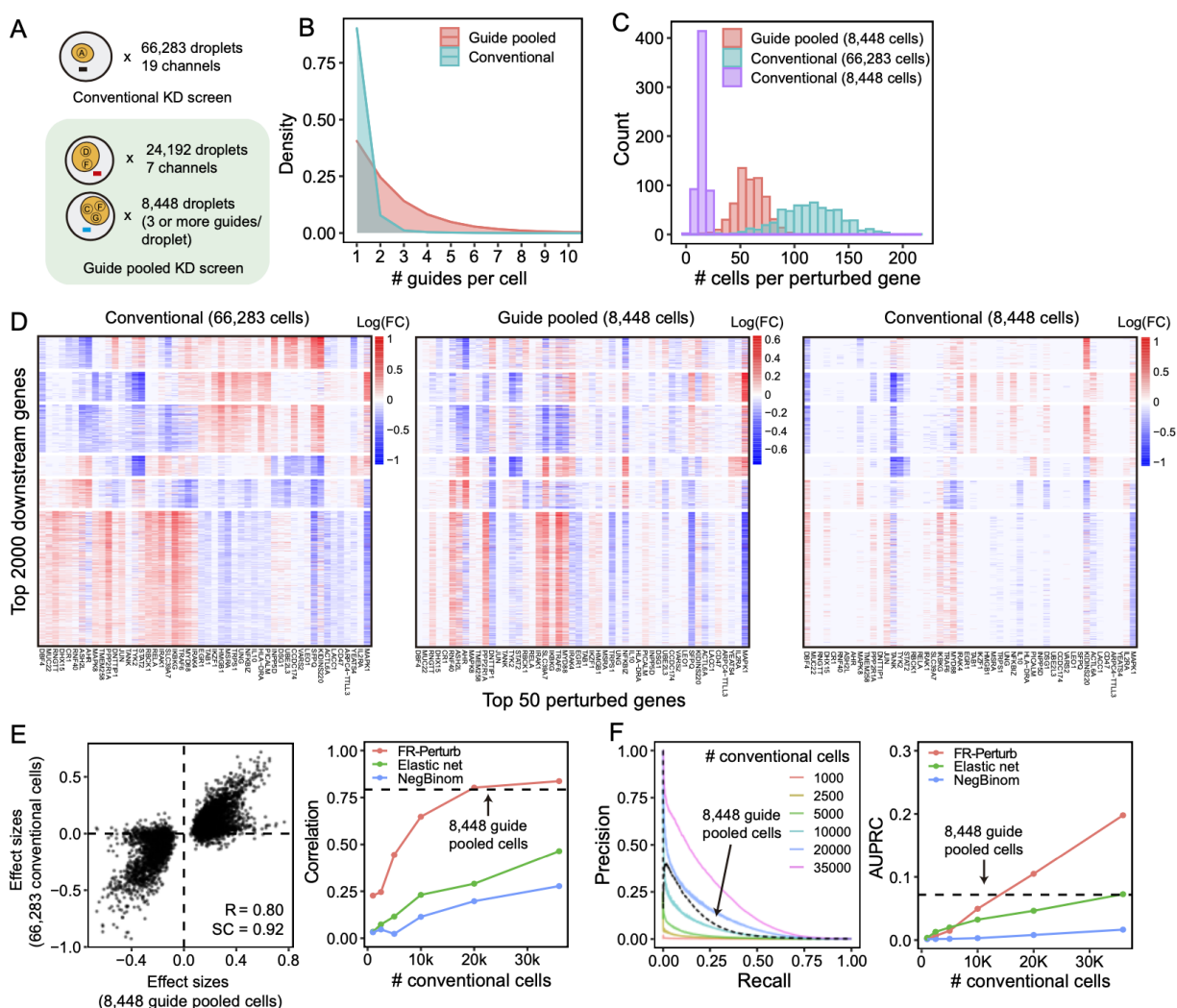
19,909) from the cell-pooled screen (X-axis) versus the same effects in the conventional screen (Y-axis). Effects are represented in terms of log-fold changes in expression relative to control cells. R, Pearson's correlation. SC, sign concordance. (Right) Held-out validation accuracy of down-sampled conventional screen and cell-pooled screen. X-axis: Size of down-sampled conventional screen by number of channels. Y-axis: Validation accuracy (in terms of correlation of effects with held-out validation dataset) of top 19,909 effects from the down-sampled conventional screen. Effects are estimated using FR-Perturb, elastic net regression, or negative binomial regression. The same method is used to estimate effects in both the down-sampled conventional data and validation data. The results from the cell-pooled screen (dotted line) are shown for estimation with FR-Perturb only (see **Figure S2D** for results using other methods). (F) (Left) Precision-recall curves computed from down-sampled conventional screen and cell-pooled screen (dotted line). True positives are the 79,100 significant effects from the held-out validation dataset, while the classification threshold being varied (X-axis) is the significance of the effects in the down-sampled conventional screen. All effects displayed are learned using FR-Perturb. (Right) AUPRCs (Y-axis) computed from the down-sampled conventional experiment when varying the number of channels (X-axis). Effects are computed using all three inference methods.

### Guide pooling yields accurate perturbation effects at high efficiency

Guide-pooled Perturb-seq was also concordant with its conventional counterpart, based on a similar evaluation scheme as above. For the guide-pooled screen, we focused on the 8,448 cells with 3 or more guides. This number of guides per cell can be achieved with sequential transduction, as done for 2 of the 7 channels (**Methods, Figure S1B**). We learned perturbation effects from both screens using FR-Perturb, with slight modifications to account for differences in the guide-pooled vs. cell-pooled screens (**Methods, Table S2**). The 5,836 significant effects from the guide-pooled cells were strongly correlated with the same effects from the conventional screen (Pearson's  $R = 0.80$ , sign concordance = 0.92) (**Figure 4E**). Thus, even if some non-linear effects exist between guides, the overall assumption of additivity holds broadly enough to infer many accurate effects. As with the cell-pooled screen, the total number of significant effects was much lower in the 8,448 guide-pooled cells vs. the full conventional screen (5,836 vs. 95,526;  $q\text{-value} < 0.05$ ), but this is expected because our conventional screen was by design much larger and more highly powered overall to enable down-sampling analyses.

The guide-pooled screen was substantially more efficient than the conventional screen per experimental unit (cell), and FR-Perturb provided more accurate effect sizes than established methods. Around 2.5x more conventionally studied cells were needed to obtain the same number of significant effects as guide-pooled cells (**Figure S2E**). Globally, the effect size patterns learned from the same number of cells (8,448 cells) were much less noisy in the guide-pooled screen than in the conventional screen (**Figure 4D**). Approximately twice as many conventional cells were required to learn effect sizes at the same correlation (**Figure 4E**) or to attain the same AUPRC (**Figure 4F**) as guide-pooled cells when comparing to a held-out validation set. This relative efficiency gain was consistent when varying the number of compared effects (**Figure S2G**). Moreover, the effect sizes inferred by FR-Perturb had substantially better validation accuracy than those from the two established inference methods in both the guide-pooled and conventional data (**Figure 4E,F, Figure S2H**). We found that around 3,200 guide-pooled cells analyzed with FR-Perturb achieved the same validation accuracy as 36,000 conventional cells analyzed with existing

approaches (**Figure S2F**), leading to an approximately 10-fold cell count and cost reduction over existing experimental and computational approaches (**Discussion**).



**Figure 4. Evaluating guide-pooled Perturb-seq versus conventional Perturb-seq.** (A) Number of channels and droplets from the conventional validation screen (top) and guide-pooled screen (bottom). We focus our analysis on the subset of 8,448 droplets from the guide-pooled screen with at least 3 guides per droplet. (B) Distribution of cells based on # of guides they contain for the full guide-pooled and conventional screens. In practice, we only directly measure the # of guides/droplet rather than guides/cell, but these quantities are equivalent given 1 cell/droplet. (C-F) See captions for **Figure 3C-F**. These analyses were conducted in an identical fashion, with the only difference that the screens are down-sampled based on cell count rather than channel count.

### Guide pooling is the more impactful compression approach

We next considered the strengths and limitations of the two compressed designs relative to conventional designs, and to each other.

Interestingly, cell pooling performed worse than the conventional approach on a per-droplet basis in both held-out real data and simulations (**Methods, Figure S3A**), but we still observed overall efficiency gains from cell pooling (**Figure 3D-F**) because it provided 7-fold more non-empty droplets per channel (rather than providing more cells per droplet per se). Thus, the reduced per-droplet power (which occurs due to dampening of perturbation effect sizes in a manner proportional to the number of cells in a droplet, **Supplementary Note**) is balanced by increased per-channel power from obtaining more non-empty droplets. We cannot confidently predict the degree of cell pooling needed to achieve the optimal balance between these counteracting effects. For plate-based assays such as sci-Plex<sup>29</sup>, our results suggest that pooling cells will decrease per-sample performance without a counteracting gain from obtaining more samples.

Conversely, the efficiency of the guide-pooled (high MOI) design readily scales with the number of guides per cell, with the best performance attained by cells with 4 or more guides in both held-out real data and simulations (**Figure S3B**). Thus, a higher number of guides per cell (achievable via sequential infections or the use of Cas12 to knock out/down multiple genes with a single construct<sup>30,31</sup>) should further improve efficiency, regardless of whether the assay is droplet-based or plate-based.

These observations have important implications for existing Perturb-seq screens, each of which already has some overloaded droplets (cell pooling) and multi-guide expressing cells (guide pooling) by chance or by design<sup>1,2,13</sup>. While these cells/droplets are often discarded, our results suggest that these cells/droplets can contain even more signal than the single-guide/single-cell containing ones and thus should be retained. To illustrate this, we used FR-Perturb to analyze a Perturb-seq knock-out screen of 1,130 genes in mouse bone marrow derived dendritic cells (BMDCs) (Geiger-Schuller et al.). In this screen, 519,535 droplets containing a single cell were obtained, of which 33% contained more than one guide by chance. By stratifying cells by the number of guides and comparing the learned effect sizes from FR-Perturb with a held-out validation subset of the data with single guide perturbations, we show that the accuracy of the effect sizes scales with the number of guides per cell and is highest in cells containing three guides (**Figure S3C**). Thus, by retaining all cells with more than one guide, the sample size of the experiment could effectively be doubled compared to the conventional approach that discards these cells (**Figure S3C**).

### **Co-functional modules and co-regulated gene programs in the LPS response**

We next leveraged the overall concordance of all perturbation data (conventional and compressed, knock-out (KO) and knock-down (KD)) to investigate the underlying regulatory circuitry of the LPS response. To maximize power, we merged droplets from the compressed and conventional screens together, then re-estimated all effects. There were 251,792 significant effects in the combined conventional and cell-pooled KO screen (131,161 effects in the combined conventional and guide-pooled KD), an increase of 16% (KD: 37%) over the conventional screen alone.

Overall, the KO and KD screens were concordant, with most of the significant effects (FDR  $q$ -value  $< 0.05$ ) attributed to relatively few (~5%) of the perturbations, each with widespread effects on many genes (**Figure 5A**). As expected, there were substantially more significant effects in the KO compared to the KD screen (251,792 vs. 131,161 effects), consistent with larger effects of KO

on the target gene's activity<sup>32</sup>. Effects significant in both screens ( $N = 26,362$ ) were highly correlated between the screens ( $R = 0.92$ ; sign consistency = 0.99; **Figure S4A-D**). The perturbations did not lead to new global cell states, such that profiles from perturbed (one or more targeting guides) and unperturbed (control guide) cells spanned the same low-dimensional space (**Figure 5C**). Thus, while many perturbations had significant and widespread effects, they did not yield radically altered phenotypic states, consistent with previous studies of this cellular response<sup>2</sup>.

We organized the perturbations and genes by clustering their effect size profiles (**Methods**), observing four broad co-regulated programs of downstream genes with correlated responses across the perturbations, and three broad co-functional modules of perturbations with correlated effects on downstream genes (**Figure 5D**).

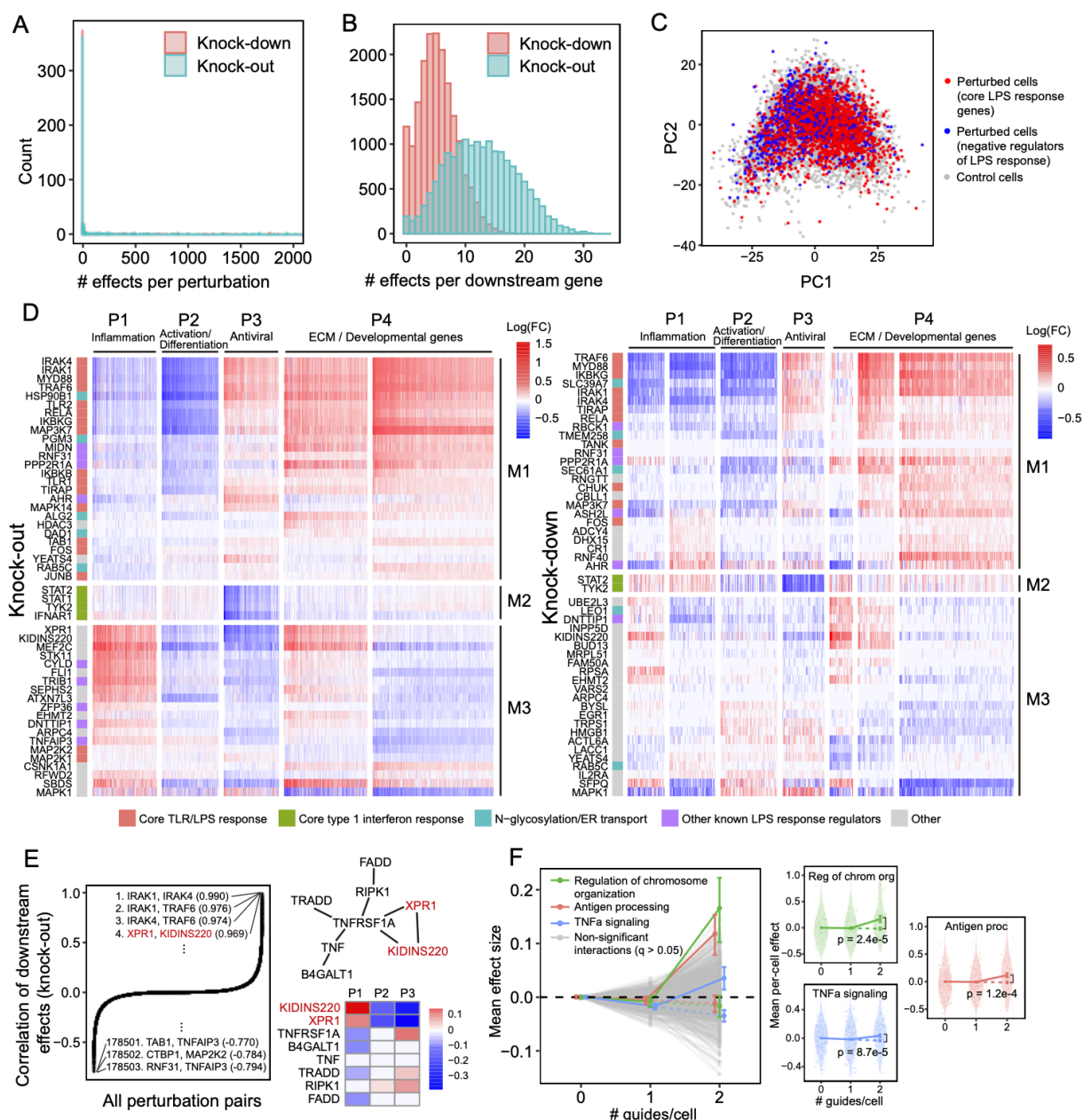
The four major co-regulated programs were present in both the KO and KD screens (**Figure 5D**), spanning key aspects of the response to LPS: inflammation (P1; cytokine, chemotaxis and LPS response genes; **Figure S4E-F**); macrophage differentiation (P2; immune cell activation, differentiation, and cell adhesion genes); antiviral response (P3; type I interferon response genes); and ECM and developmental genes (P4) (**Table S3**). Inflammation (P1) and the antiviral response (P3) are known to be regulated by LPS signaling through AP1/NF- $\kappa$ B and IRF3, respectively<sup>33</sup>, and were mostly anti-correlated in their responses to perturbation in our screen, consistent with reports that downregulation of the inflammatory response can lead to upregulation of type I interferon response<sup>34,35</sup>. Inflammatory signaling is known to lead to macrophage differentiation<sup>36</sup>, but almost all perturbations with significant effects on inflammation (P1) (in any direction) down-regulated macrophage differentiation (P2). This suggests that additional factors beyond inflammatory signaling mediate macrophage differentiation in response to LPS<sup>37</sup>.

Of the three major co-functional modules, KO/KD of the first module (M1) resulted in strong down-regulation of inflammation and macrophage differentiation (P1-2) and upregulation of the antiviral response and ECM/developmental genes (P3-4) (**Figure 5D**). M1 was mainly composed of core TLR/LPS response genes and genes directly up- or downstream of the pathway<sup>33</sup>, including MYD88, IRAK1, IRAK4, RELA, TRAF6, TIRAP, IKBKB, IKBKG, TAB1, TANK, TLR1, TLR2, MAPK14, MAP3K7, FOS, JUNB, and CHUK. Given the known function of these genes, we expect that their KO/KD will lead to down-regulation of inflammation and macrophage differentiation (P1-2), as we indeed observed. Other genes in M1 previously shown to down-regulate TNF and the inflammatory response when knocked out<sup>25</sup> included two LUBAC complex proteins (RBCK1 and RNF31), genes in the OST complex (DAD1, TMEM258) and ER transport (HSP90B1, SEC61A1, ALG2), and other genes with diverse functions (MIDN, AHR, PPP2R1A, ASH2L). M1 also included two additional ER transport genes not previously implicated in immune pathways (RAB5C, PGM3), highlighting the important role of N-glycosylation and trafficking in macrophage activation<sup>38</sup>.

KO/KD of the second co-functional module (M2) primarily resulted in strong downregulation of the antiviral program (P3), with weak/mixed effects on other programs. M2 comprised four genes known to be core components of the type 1 interferon response<sup>39</sup> – STAT1, STAT2, TYK2, and IFNAR1 – for which downregulation of the antiviral program in response to their perturbation is expected.

KO/KD of the third and final co-functional module (M3) resulted in upregulation of inflammation (P1), downregulation of macrophage differentiation and the antiviral response (P2-3), and mixed effects on ECM/development (P4). M3 included many genes with known inhibitory effects on inflammation, including ZFP36, an RNA-binding protein that destabilizes TNF mRNA<sup>40</sup>, enzymes CYLD and TNFAIP3 involved in deubiquitination of NF- $\kappa$ B pathway proteins<sup>41,42</sup>, pseudokinase TRIB1 and ubiquitin ligase RFWD2 which are involved in degradation of JUN<sup>43,44</sup>, and RELA-homolog DNTTIP1<sup>25</sup>. Other genes in M3 included transcription factors (MEF2C, FLI, and EGR1), chromatin modifiers (EHMT2, ATXN7L3), and kinases (CSNK1A1, STK11).

Interestingly, two of the M3 genes with particularly strong effects on all programs did not have prior immune annotations – XPR1, a retrovirus receptor involved in phosphate export, and KIDINS220, a transmembrane scaffold protein previously reported in neurons<sup>45</sup>. In the KO screen, this pair of genes had the fourth highest correlation of downstream effects ( $R = 0.97$ ) among all  $\binom{598}{2} = 178,503$  perturbation pairs (**Figure 5E**), following IRAK1/IRAK4, IRAK1/TRAF6, and IRAK4/TRAF6 which are all known to form a physical LPS signaling complex<sup>33</sup>. In AP-MS data<sup>46</sup>, XPR1 and KIDINS220 physically associate with each other and TNF receptor TNFRSF1A. Knockout of TNFRSF1A in our screen results in effects opposite to XPR1/KIDINS220 KO (**Figure 5E**), suggesting a possible inhibitory effect of this complex on TNFRSF1A.



**Figure 5. Analysis of knock-out and knock-down perturbation effects in the LPS response.**

(A) Distribution of perturbed genes based on their number of significant effects ( $q < 0.05$ ) on downstream genes. (B) Distribution of downstream genes based on how many perturbed genes significantly affect their expression. (C) PCA of perturbed and unperturbed cells based on the expression of the top 2,000 most variable genes. Grey points represent cells containing a non-targeting guide only. Red points represent cells containing a guide for genes involved in the core LPS/TLR response (IKBKB, IKBKG, IRAK1, IRAK4, MAP2K1, MAP3K7, MAPK14, MYD88, RELA, TIRAP, TLR1, TLR2, and TRAF6). Blue points represent cells containing a guide for genes known to be negative regulators of the LPS response (CISH, CYLD, STAT3, TNFAIP3, TRIB1, and ZFP36). (D) Heatmaps of perturbation effect sizes (inferred with FR-Perturb) from the knock-out (left) and knock-down (right) screens. Rows: top 50 perturbed genes

based on their average magnitude of effects on all downstream genes. Columns: top 2,000 downstream genes based on the average magnitude of effects of all perturbed genes acting on them. Rows and columns are clustered using Leiden clustering. Clusters are labelled based on their GO enrichment terms. All effects with  $q > 0.2$  are whited out. **(E)** (Left) Correlation of knock-out effect sizes (y-axis) between all pairs of perturbed genes (x-axis). Top and bottom gene pairs are labelled. (Top right) Graph of all perturbed genes that physically interact with XPR1 and/or KIDINS220, based on AP-MS data from Bioplex 3.0<sup>46</sup>. Edges represent physical interaction. (Bottom right) Mean effects of perturbed genes from top right on P1-P3. **(F)** Analysis of genetic interaction effects. (Left) Effect sizes relative to control (y-axis) of cells containing 0, 1, or 2 guides (x-axis) within each perturbation module (lines connecting three dots). Modules with significant effects ( $q < 0.05$ ) are highlighted in color and labeled, with the expected effect of cells containing two guides in the module represented with a dotted line. All other non-significant modules are represented in grey. Error bars represent standard errors of effect sizes obtained from bootstrapping. (Right plots) Violin plots of the mean effects of individual cells (relative to control) containing 0, 1, or 2 guides in the three perturbation modules with significant interaction effects. Dotted line represents the expected effect relative to control of cells with 2 guides. P-values are computed from permutation testing.

### Guide pooling reveals second-order genetic interactions

Genetic interactions (non-additive effects) between two or more genes can in principle be inferred from cells containing two or more guides, which are generated by chance when transducing cells at low or high MOI. **(Figure 4B)**. Here, guide-pooling can provide increased efficiency compared to the conventional approach, like in the first-order case **(Supplementary Note)**.

We first attempted to estimate second-order interaction effects and their p-values from the guide-pooled screen and corresponding conventional KD screen by adding interaction terms to the perturbation design matrix **(Methods)**. However, although we could generate point estimates of second-order effects<sup>2</sup>, none of these effects was significant in either screen due to insufficient power **(Figure S5A)**, even with a lax significance threshold ( $q < 0.5$ ).

To increase power, we aggregated perturbations into modules defined by GO annotations **(Table S4A)** and learned the overall impact of second-order interactions within and between each module on each gene program **(Methods)**. Here, we define an interaction effect as the deviation from the sum of first-order effects for cells that contain any two perturbations from either the same module (intra-module interactions) or two different modules (inter-module interactions) **(Methods)**. To ensure adequately sized groupings, we aggregated perturbations into 490 (possibly overlapping) modules each with at least 20 genes, such that any pair of perturbations in each module was represented in an average of 87 cells in the guide-pooled screen (conventional: 30 cells) **(Figure S5B)**. We also constructed 30 non-overlapping modules by clustering the original 490 modules **(Methods)**, resulting in  $\binom{30}{2} = 435$  module pairs among which we could compute inter-module interactions. To increase power, we grouped downstream genes by their program (P1-4) membership **(Figure 5D)**, computing mean effects on these four programs rather than on individual genes. The results from this analysis represent the extent of intra- and inter-module interactions on each key program.

We detected three co-functional modules with significant ( $q < 0.05$ ) intra-module interaction effects on at least one program from the guide-pooled screen (**Figure 5F; Table S4B**), while we detected no significant interactions from the substantially larger conventional screen (even at  $q < 0.5$ ) (**Figure S5C; Table S4C**). Two of the significant interaction effects – with genes for regulation of chromosome organization ( $p = 2.4 \times 10^{-5}$ ) and antigen processing ( $p = 1.2 \times 10^{-4}$ ) – had insignificant first-order effects on the antiviral program (P3), while having significant positive second-order effects. The third, TNFa signaling, had a significant negative first-order effect on the inflammatory/LPS program (P1) ( $p = 2.0 \times 10^{-4}$ ) and significant positive second-order effect ( $p = 8.7 \times 10^{-5}$ ). This effect is consistent with the reported non-linear relationship between gene dosage and TNF signaling activity when comparing heterozygous versus homozygous KO mice for either TNF<sup>47</sup> or the TNF receptor TNFRSF1A<sup>48</sup>. Interestingly, we did not observe any significant inter-module interactions from either screen (**Figure S5D; Table S4D,E**), which may suggest that perturbations in different modules are less likely to interact with each other<sup>49,50</sup>.

### Integration of perturbation effects with genome-wide association studies implicates disease-relevant perturbations

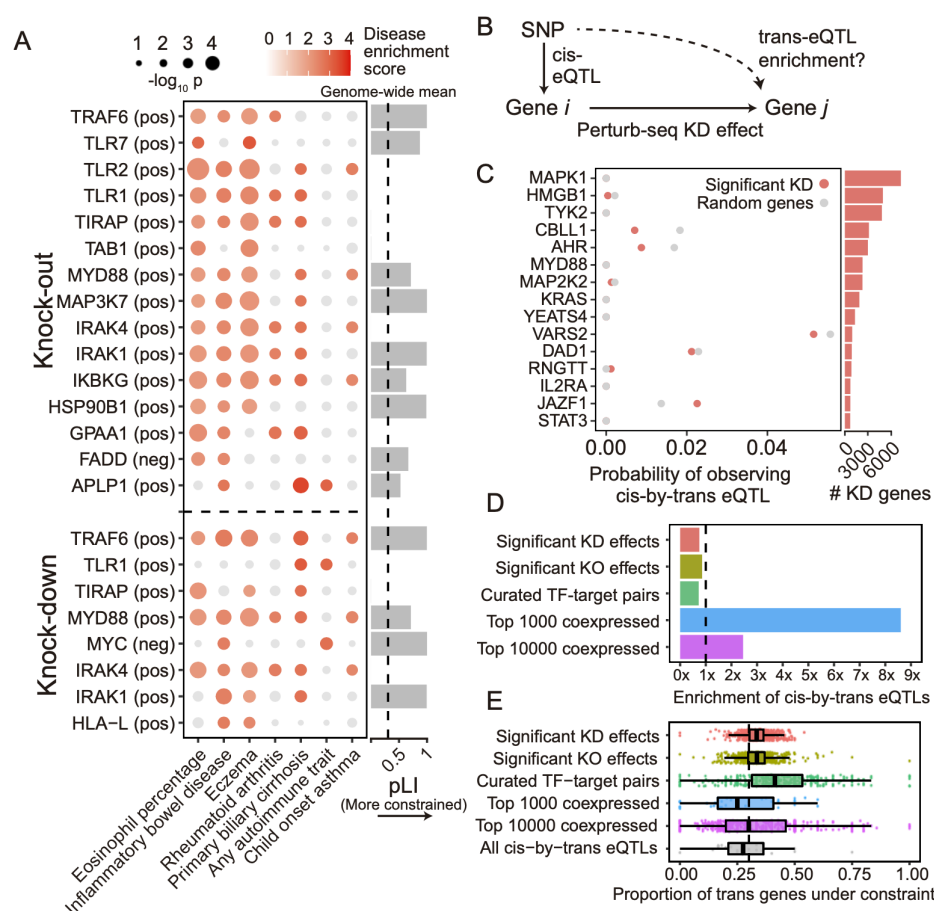
Because dysregulation of innate immune responses plays a key role in many human diseases<sup>51</sup>, we next asked whether the perturbation effects learned from our *in vitro* screens can help identify disease-relevant genes and processes. *In vitro* screens may be especially helpful for this aim given that many of the perturbed genes from our screens are under strong selective constraint in human populations (**Figure S6A**), making them challenging to directly connect to disease through genome-wide association studies<sup>52</sup> (GWAS) due to fewer common variants in or around the gene<sup>53,54</sup>. To investigate this, we obtained summary statistics from GWAS of 64 distinct human diseases and traits (**Table S5A**), including autoimmune diseases and blood traits, as well as non-immune traits/diseases (e.g. height, BMI, schizophrenia, type 2 diabetes). Using sc-linker<sup>55</sup>, we computed the overall heritability enrichment of these 64 traits/diseases in SNPs in/around genes comprising perturbation modules M1-3 (**Methods**). We observed significant heritability enrichment ( $p < 0.001$ ) for M3 (genes that suppress the LPS response) for two blood traits (lymphocyte and neutrophil percentage), but did not observe significant enrichment for M1 (positive regulators of the LPS response) or M2 (genes involved in the antiviral response) for any traits (**Figure S6B**).

Instead, we hypothesized that if a perturbed gene is important for disease, then disease heritability may be enriched near the *downstream genes* it affects<sup>12,56</sup>. To test this hypothesis, we constructed two “perturbation signatures” for each perturbed gene that include all genes that are significantly upregulated (“negative” targets) or downregulated (“positive” targets) by its KO/KD. We retained signatures with at least 100 genes, resulting in a total of 1,634 perturbation signatures from both the KO and KD screens. We also constructed signatures corresponding to the gene programs P1-4 (**Figure 5D**). As above, we used sc-linker to test for disease heritability enrichment for each signature/phenotype pair (**Methods**).

23 signatures associated with 16 perturbed genes had significant heritability enrichment scores for at least two phenotypes ( $p$ -value  $< 0.001$ ). Meanwhile, 7 phenotypes that all reflect immune or blood traits (inflammatory bowel disease, eczema, rheumatoid arthritis, asthma, primary biliary cirrhosis, and eosinophil percentage) had significant scores for at least two perturbation signatures (**Figure 6A; Figure S6C,D; Table S5B,C**). Most of the significant signatures (15/23) were from

the KO screen, suggesting that the expression effects from KO are more suited for this analysis (either because they are more disease-relevant or more powered due to capturing more effects). Among the downstream programs P1-4, we observed significant enrichment from only P2 on three immune traits: inflammatory bowel disease, eczema, and primary biliary cirrhosis (**Figure S6B**).

Most of the significant signatures (17/23) were from genes in core LPS and TLR signaling pathways that fall into perturbation module M1 (even though M1 did not exhibit any direct heritability enrichment itself; **Figure S6B**): TRAF6 (positive), TLR7 (positive), TLR2 (positive), TLR1 (positive), TIRAP (positive), TAB1 (positive), MYD88 (positive), MAP3K7 (positive), IRAK4 (positive), IRAK1 (positive), and IKBKG (positive). Other significant signatures include HSP90B1 (positive), an ER transport gene important for innate immunity<sup>57</sup> that is co-functional with the core LPS genes (**Figure 5D**); FADD (negative), a pro-apoptotic gene downstream of LPS signaling that serves for negative feedback<sup>33</sup>; MYC (negative), an oncogene with known immunosuppressive effects<sup>58,59</sup>; and poorly characterized pseudogene HLA-L. The two remaining significant signatures are for genes whose functions are not previously associated with the immune system, including APLP1 (an amyloid beta precursor-like gene primarily involved in brain function that interestingly contains a missense variant associated with severe influenza<sup>60</sup>) and GPAA1 (involved in anchoring proteins to the cell membrane). Thus, by leveraging gene-gene links learned from our screens, we were able to identify disease-relevant genes that we were underpowered to detect through direct heritability analyses (**Discussion**).



**Figure 6. Integration of population-genetic screens with Perturb-seq.** (A) Heritability enrichment scores of signatures comprising genes significantly modulated by perturbations (rows) across human traits (columns), computed using sc-linker<sup>55</sup>. “Pos” indicates the set of genes whose expression changes in the same direction as the perturbed gene (*i.e.*, downregulated by the perturbation), with the opposite applying to “neg”. Displayed are all perturbation signatures and traits with at least 2 significant ( $p < 0.001$ ) effects. Non-significant scores are greyed out. (Barplot) Probability of loss-of-function intolerance<sup>54</sup> (pLI) of the corresponding perturbed gene. (B) Schematic of eQTL integration analysis, aiming to test whether *trans*-regulatory relationships learned from Perturb-seq are also present in eQTL studies. For all gene pairs in which gene *i* exerts an effect on gene *j* (*i.e.*, has a significant knock-down effect in our Perturb-seq screen), we would expect that gene *i* and gene *j* are enriched for *cis*-by-*trans* eQTLs. (C) Using data from an eQTL study closely matching our cell type and treatment<sup>26</sup>, shown is the probability of observing significant *cis*-by-*trans* eQTLs among the top 15 perturbed genes from our knock-down screen and their affected downstream genes (red) compared to random downstream genes (grey). (D) Enrichment of significant *cis*-by-*trans* eQTLs among various sources of gene-gene pairs: significant KO/KD effects (representing significant gene-gene effects from our KO and KD screens, respectively), curated transcription factor (TF) and target gene pairs<sup>61</sup>, and the top 1,000/10,000 most co-expressed gene pairs (based on correlation of expression across samples) from the eQTL dataset. Enrichment is computed relative to random *trans* genes for each *cis* gene, then averaged over all *cis* genes. (E) Selective constraint on *trans* genes from (D) plus all significant *cis*-by-*trans* eQTLs from the Fairfax dataset. Each point represents a *cis* gene, while the x-axis represents the proportion of the *trans* genes for each *cis* gene that are under selective constraint (determined as having a pLI > 0.5).

### **Perturbation effects do not overlap with *trans*-regulatory links from an eQTL study in primary monocytes**

*Trans*-genetic gene regulation (*i.e.*, regulation of gene expression distal to the given SNP) has been proposed as a primary mediator of genetic effects on human disease<sup>62</sup>. *Trans*-genetic gene regulation can be studied through either population-level genetic data (via expression quantitative trait loci (eQTL) studies<sup>63,64</sup>), or through experimental perturbation of gene expression<sup>12</sup>, such as the screens conducted in our study. Although both types of data can in principle be used to learn the same *trans* effects, their consistency with each other has not been empirically evaluated.

We therefore compared gene-gene regulatory links between our Perturb-seq screen and a *trans*-eQTL analysis in primary patient-derived monocytes treated with LPS<sup>26</sup> ( $N = 432$ ), closely matching our cell line. We define a gene-gene regulatory link in eQTL studies based on *cis*-by-*trans* colocalization, where a *cis*-eQTL for gene *i* is also a *trans*-eQTL for gene *j* via a (presumed) *trans*-regulatory effect of gene *i* on gene *j* (**Figure 6B**). Here, we assume that a perturbation of a *cis*-eQTL on the expression of gene *i* is analogous to the experimental KD in our system. We used coloc<sup>65</sup> to compute the posterior probability of *cis*-by-*trans* colocalization, while accounting for linkage disequilibrium between SNPs (**Methods**). To determine whether the regulatory links learned for a given perturbed gene *i* from Perturb-seq are reflected in the eQTL analysis, we compared the proportion of downstream genes *j* of gene *i* in Perturb-seq that colocalize with gene *i* in the eQTL study,  $P(\text{coloc}_{\text{gene } i \rightarrow \text{gene } j})$ , with the proportion of random expressed genes that colocalize with *i*,  $P(\text{coloc}_{\text{gene } i \rightarrow \text{random gene}})$  (**Methods**).

Surprisingly,  $P(\text{coloc}_{\text{gene } i \rightarrow \text{gene } j})$  was slightly lower than  $P(\text{coloc}_{\text{gene } i \rightarrow \text{random gene}})$  for individual perturbed genes  $i$  (**Figure 6C**), as well as when aggregating across all perturbed genes (**Figure 6D**). Moreover, we observed no relationship between either the significance or magnitude of the effect of gene  $i$  on gene  $j$  and  $P(\text{coloc}_{\text{gene } i \rightarrow \text{gene } j})$  (**Figure S7A**). We observed similar negative results when obtaining gene-gene links from our KO data or from a curated list of transcription factor-target gene pairs<sup>61</sup> (**Figure 6D**). Using an alternative way of quantifying gene-gene links in eQTL studies that does not make assumptions about the number of causal variants (*i.e.*, bivariate Haseman-Elston regression to estimate genetic correlation of expression<sup>66</sup>; **Methods**) yielded similar results (**Figure S7B,C**).

Conversely, we did observe significant enrichment of *cis*-by-*trans* eQTLs in gene pairs co-expressed in the same eQTL study (**Figure 6D**), as has been observed in other trans-eQTL studies<sup>63</sup>. Notably, co-expression in eQTL datasets is dominated by environmental effects rather than genetic effects<sup>67</sup>. Thus, given that the two effects are independent across samples, we would not ordinarily expect the most strongly co-expressed genes to be enriched for *cis*-by-*trans* eQTLs, suggesting that they may be confounded in part by unmodelled technical artifacts or inter-cellular heterogeneity (**Supplementary Note**). We also observed that the level of negative selection on the *trans* gene mirrored the patterns of *cis*-by-*trans* eQTL enrichment (or lack thereof) we observed in the previous analyses (**Figure 6E**), suggesting that our power to detect *cis*-by-*trans* eQTLs was affected by selection-induced depletion of SNPs affecting the *trans* genes<sup>54,68</sup> (**Supplementary Note**).

## DISCUSSION

Here, we evaluated a new approach for conducting Perturb-seq based on generating composite samples, which involves either overloading microfluidics chips to generate droplets containing multiple cells (cell pooling), or infecting cells at high MOI so that each cell contains multiple guides (guide pooling). We also propose a new method, FR-Perturb, to estimate perturbation effect sizes from composite samples, which increases power by estimating sparsity-constrained effects on latent gene expression factors rather than on individual genes. We tested our approach by perturbing 598 immune-related genes in a human macrophage cell line. We found that our experimental approaches of cell pooling and guide pooling, combined with the use of FR-Perturb to infer effect sizes, lead to large cost reductions over conventional Perturb-seq while maintaining the same accuracy. Guide pooling also significantly increases power to detect genetic interaction effects and reduces the number of cells needed for screening.

Inference with FR-Perturb leads to substantially improved out-of-sample validation accuracy over conventional gene-by-gene methods (*e.g.*, elastic net, negative binomial regression) in both conventionally generated data and compressed data. FR-Perturb is thus useful for inferring effects in any type of Perturb-seq screen, even conventional screens that do not adopt our proposed experimental changes. The improved performance of FR-Perturb in both conventional and compressed settings likely stems from perturbation effect sizes being inferred on latent gene expression factors that aggregate many co-expressed genes, thereby denoising the expression counts of individual genes which are especially noisy/sparse in single-cell data.

Compressed Perturb-seq using cell-pooling led to a 4 to 20-fold cost reduction over existing approaches in our experiments, while guide-pooling led to a 10-fold cost reduction, based on

matching the number of samples needed to obtain equivalent out-of-sample validation accuracy. These numbers factor in both our experimental changes (cell-pooling or guide-pooling) combined with the use of FR-Perturb to infer effects, compared to existing experimental approaches combined with the use of existing methods to infer effects (elastic net or negative binomial regression). For cell-pooling, we account for the fact that cell-pooled channels require deeper sequencing than conventional channels (4x per channel in our case): assuming that one cell-pooled channel is equivalent to 10 conventional channels and costs are equally divided between library preparation and sequencing in the conventional screen (as was the case in our screen), this implies a relative cost of  $0.5 * (0.1x \text{ library preparation costs} + 4 * 0.1x \text{ sequencing costs}) = 0.25x$ , or a 4x reduction, of cell pooling. We also provide an estimate based on our analysis that only 1/5 of a cell-pooled channel is sufficient to achieve the same validation accuracy as 10 conventional channels, resulting in a relative cost of  $0.5 * (0.2 * 0.1x \text{ library preparation costs} + 0.2 * 4 * 0.1x \text{ sequencing costs}) = 0.05x$ , or a 20x reduction. Because we did not generate enough conventional data to match the performance of a single cell-pooled channel (which cannot be sub-divided in a real experiment), we report the cost reduction from cell pooling as a range rather than an exact estimate. On the other hand, guide pooling produces composite samples at the level of cells and does not require increased per-channel sequencing costs, so the cost reduction can be directly computed as the ratio of required conventional cells to guide-pooled cells (10x in our case).

Compressed Perturb-seq reduces costs due to RNA library preparation without altering the sequencing step of scRNA-seq. Thus, it can in principle be paired with approaches that increase the efficiency of sequencing via new technologies<sup>69</sup> or targeted sequencing<sup>70</sup>, resulting in further improvements to the efficiency of Perturb-seq. Concurrent results also demonstrate the power of compressed screening with bio-chemical perturbations in high-fidelity cellular model systems (Mead et al., companion manuscript).

Cell pooling and guide pooling are complementary approaches with different strengths and limitations. Unlike cell pooling, guide pooling has the drawbacks that it requires that nonlinear interaction effects do not systematically bias phenotypes, and it potentially suffers from cellular toxicity caused by multiple viruses infecting each cell and/or multiple double stranded breaks. Meanwhile, unlike guide pooling, cell pooling has the drawbacks that it requires increased sequencing depth per channel to account for more non-empty droplets, and it loses per-droplet signal due to dilution of effect sizes (**Supplementary Note**). Due to the latter fact, cell pooling requires many more cells than guide pooling to achieve the same performance, which can be prohibitive in certain settings where cell count is limited<sup>8,13</sup>. Because guide pooling performs best with high guide number per cell (4 or more), whereas cell pooling does not perform well with high cell count per droplet, we posit that guide pooling (but not cell pooling) can be readily scaled up to very compressed designs (in which case the use of knock-down over knock-out and Cas12 over Cas9 may be desirable to avoid cellular toxicity), likely leading to even larger efficiency gains than we observed in our screens.

An additional key advantage of guide pooling over cell pooling is that guide pooling naturally allows for the study of higher-order interaction effects. In our study, we were underpowered (even with guide pooling) to detect second-order interaction effects between individual gene pairs. However, we detected significant intra-module interaction effects from the guide-pooled but not conventional screen, serving as a proof-of-concept that such signal can be detected in the guide-

pooled screen, and may be further probed in more powered future experiments. The efficiency gains brought about from guide pooling can in theory counteract the exponential growth of gene combinations (given that various assumptions are satisfied), potentially making it the only tractable way to systematically study higher-order interaction effects (**Supplementary Note**).

By integrating data from GWAS, our screens highlighted perturbed genes with downstream genes enriched for disease heritability. Many of these perturbed genes are under strong selective constraint and would require up to millions of samples to detect in GWAS<sup>71</sup>. Thus, our analysis represents a potential way to circumvent the issue of negative selection removing GWAS signal from large-effect disease-relevant genes, a key challenge for biological interpretation of common-variant GWAS.

Gene-gene effects learned from our Perturb-seq screens were not enriched for *cis*-by-*trans* eQTLs in a closely matched cell type and treatment. Many possible explanations exist for this observation, including (1) insufficient power to detect *trans*-eQTLs in the eQTL dataset, (2) biological differences between our cell line and primary monocytes used in the eQTL study, (3) large differences in the magnitude of perturbation between experimental KO/KD and eQTLs, and (4) confounders in the eQTL dataset (**Supplementary Note**). Explanation (1) can in theory be addressed with larger *trans*-eQTL studies<sup>63</sup>, though such studies often suffer from issues with confounding/intercellular heterogeneity, as evidenced by very low reported out-of-sample replication accuracy and substantial overlap (>50%) of detected *trans*-eQTLs with variants known to influence cell type proportion<sup>63</sup>. Meanwhile, single-cell eQTL studies<sup>72</sup> can potentially address explanation (4), though such studies suffer from low power relative to sample size (~1,000 significant *trans*-eQTL effects detected from ~1.2 million cells<sup>72</sup> versus ~200,000 *trans* perturbation effects detected from ~100,000 cells in our screen). We propose that our compressed screen is a powerful tool to learn *trans*-effects on gene expression, while additional work is needed to fully reconcile the differences between population-level genetic screens and experimental perturbation screens.

### Author contributions

BC and AR conceived of the project and designed the experiments. LB, JB, BS, JF, and BC ran the experiments with input from AR. DY and BC implemented FR-Perturb and conducted computational and biological analyses with input from AG. CF and KD assisted with computational analyses. KGS and BE provided data for the mouse BMDC Perturb-seq. DY, AG, AR, and BC wrote the paper with input from all authors.

### Acknowledgements

We thank Atray Dixit for early discussions on efficient screens and Orit Rozenblatt-Rosen for discussions and help. BC was supported by the Broad Fellows program and a Merkin Institute Fellowship at the Broad Institute. DY was supported by the NSF Graduate Research Fellowship Program (Grant #1745303). AG was supported by R01 HG012133 and R01 HG006399. AR was supported by an NHGRI Center of Excellence in Genome Science grant (CEGS; RM1HG006193; AR), the Howard Hughes Medical Institute, and the Klarman Cell Observatory and Klarman Incubator at the Broad Institute. AR was a Howard Hughes Medical Institute Investigator when this study was initiated.

### **Conflict of Interest statement**

AR is a co-founder and equity holder of Celsius Therapeutics, an equity holder in Immunitas, and was a scientific advisory board member of ThermoFisher Scientific, Syros Pharmaceuticals, Neogene Therapeutics and Asimov until July 31, 2020. AR, BE, and KGS are employees of Genentech from August 1, 2020, March 10, 2022, and November 16, 2020, respectively. AR and KGS have equity in Roche. BC and AR are co-inventors on patents filed by the Broad Institute relating to Perturb-seq and compressed sensing methods of this paper.

# Methods

## EXPERIMENTAL PROCEDURES

### Cell culture and stimulation

THP-1 cells (ATCC, TIB202) were cultured in RPMI medium (ATCC, 30-2001) supplemented with 10% FBS (ATCC, 30-2020) and 0.05mM 2-mercaptoethanol (Sigma Aldrich, M7522). Cells were maintained between 0.8 and 2 million cells per milliliter.

Cell lines for knockout (KO) and knockdown (KD) screens were engineered with lentiviral vectors containing Cas9 (pxpr311) and dCas9-KRAB (pxpr121), respectively. Viruses were prepared using a previously published protocol (<https://portals.broadinstitute.org/gpp/public/dir/download?dirpath=protocols/production&filename=TRC%20shRNA%20sgRNA%20ORF%20Low%20Throughput%20Viral%20Production%20201506.pdf>) and concentrated by centrifugation in a column with a cut size of 100kDa (MilliporeSigma UFC903096). Cells were transduced by spinfection as previously described (<https://portals.broadinstitute.org/gpp/public/resources/protocols>).

THP-1 cell lines were infected with sgRNA libraries (described below) at a multiplicity of infection (MOI) specific for each guide-pooled experiment. 12 hours after spinfection, cells and media were diluted 1:10 and cells were allowed to recover for 48h. Cells were selected with puromycin (2 µg/mL) for four days. The selected cells were differentiated into macrophages by stimulation in 20ng/mL phorbol 12-myristate 13-acetate (Sigma Aldrich, P8139-1mg) for 24 hours. Cells were then allowed to rest in normal culture medium for 48 hours before stimulation in medium containing 100ng/mL LPS (MilliporeSigma, L4391-1mg) for 3 hours.

### Guide library production and validation

sgRNAs for the perturbed panel of genes (described below) were designed using the Crispr-Pick tool from the Broad Institute. Four distinct sgRNAs were designed for each perturbed gene. In addition, 500 non-targeting sgRNAs and 500 safe-targeting sgRNAs (*i.e.*, guides targeting intergenic regions of the genome) were included. Oligonucleotide libraries were synthesized by Twist Biosciences, then amplified and inserted into a CROP-Seq vector<sup>4</sup> with sgOpti scaffold (Addgene #106280) via Gibson assembly. Cloned libraries for KO, KD, and control sgRNAs (non-targeting and safe-targeting) were sequence-validated as previously described ([https://portals.broadinstitute.org/gpp/public/dir/download?dirpath=protocols/production&filename=cloning\\_of\\_oligos\\_for\\_sgRNA\\_shRNA\\_nov2019.pdf](https://portals.broadinstitute.org/gpp/public/dir/download?dirpath=protocols/production&filename=cloning_of_oligos_for_sgRNA_shRNA_nov2019.pdf)). Viral libraries were produced as described above (without concentration), and an MOI was determined by transfecting cells with scaled dilutions of the virus covering a 100-fold dynamic range and quantifying survival rate after selection.

### Conventional Perturb-Seq, cell-pooling, and guide-pooling (scRNAseq & dialout library production)

For conventional screens, the infected (MOI 0.25) and stimulated THP-1 cell suspension was prepared for droplet generation according to the manufacturer's suggested protocol (10x Genomics, CG00053 Rev C). Channels aiming to recover 5,000-10,000 cells were loaded on the 10x Chromium Controller and the protocol was followed according to the manual for Chromium Next GEM Single Cell 3' Reagent Kits v3.1 (CG000315 Rev C).

For cell-pooling (MOI 0.25), the standard 10x single cell 3' RNAseq protocol (Chromium Next GEM Single Cell 3' GEM, Library & Gel Bead Kit v3.1 PN-1000121) was run according to manufacturer's recommendation, except the concentration of cells was increased to co-encapsulate multiple cells per droplet (250,000 cells loaded per channel).

For guide-pooling, cells were infected at an MOI of 10 before selection and stimulation, or were left to rest for 2 days after initial infection before infecting a second time at an MOI of 10 before selection and stimulation (**Figure S1B**). High MOI cells were loaded into droplets as in the conventional screens.

After the generation of double-stranded cDNA, part of the whole transcriptome amplification (WTA) product was set aside for targeted amplification to recover the perturbation barcode. 10ng of WTA from each channel were input into 8 cycles of PCR (primer 1 CTACACGACGCTCTTCCGATCT; primer 2 GTGACTGGAGTTCAGACGTGTGCTCTTCCGATCTTGTGGAAAGGACGAAACACC). The sample underwent a 1x AMPure XP Reagent SPRI clean (Beckman Coulter A63881) and was amplified for another 9 cycles with 8bp indexed PCR primers and purified with a 0.7x SPRI clean (primer 1 AATGATACGGCGACCACCGAGATCTACACTCTTTCCCTACACGACGCTC, primer 2 CAAGCAGAAGACGGCATACGAGATGTCGAGCAGTGACTGGAGTTCAGACGTGTGCTCTTCCGATCT).

## COMPUTATIONAL PROCEDURES

### Selecting genes to be perturbed

A set of perturbed genes was compiled from several sources (**Table S1**). These included: a manually curated list of 35 canonical LPS response genes; the top 100 genes from a previous genome-wide CRISPR screen for regulation of TNF expression after LPS stimulation<sup>25</sup>; 100 genes identified as being a *cis* eQTL target of SNPs that were (in total) associated with *trans* eQTL effects for at least 4 downstream genes in primary monocytes treated with LPS<sup>26</sup>; 95 genes near high confidence variants in IBD GWAS loci<sup>73</sup>; 108 genes associated with Mendelian disorders identified by search for “bacterial infection” in the Online Mendelian Inheritance in Man (OMIM) database<sup>74</sup> and 115 Mendelian genes similarly identified by “NF-kappa-b” search; and 173 genes reported in studies identified by a GWAS Catalog<sup>75</sup> search for “infection” with diseases/traits related to liver disease and HIV-1 infection excluded.

The (perhaps surprisingly small) intersections between gene lists from these sources is depicted in **Figure S1A**. The final list of 598 perturbed genes was obtained by intersecting genes expressed in THP-1 cells with the combined list of 758 genes from all sources.

### Generating expression and perturbation design matrix

Starting with raw Illumina BCL files from the sequencing output, the “cellranger mkfastq” command with default parameters (from the 10x CellRanger tool v6.0.1; <https://support.10xgenomics.com/single-cell-gene-expression/software/downloads/latest>) was used to generate FASTQ files. The “cellranger count” command with default parameters was used

to align the expression reads to the GRCh38 build of the human transcriptome and generate a gene expression count matrix (see below for details on normalization of expression counts).

To generate the droplet by perturbation design matrix, paired-end reads (in FASTQ format) containing a droplet barcode and UMI on read 1 and sgRNA sequence on read 2 were aligned using Bowtie2 as follows. Read 2 reads were aligned to a reference constructed from the labeled sgRNA sequences using the --local option with default parameters, which performs local read alignment. Then, using a custom script, droplet barcodes were matched to the mapped guides for each paired-end read. A guide was called as “present” in a droplet if there were at least 5 UMIs for each droplet barcode / guide barcode pairs.

### Inference using FR-Perturb

From the sequencing output of each of our Perturb-seq experiments, two matrices were directly generated (see above):

- $N \times G$  raw gene expression count matrix  $\mathbf{Y}$ , where  $N$  is the number of droplets and  $G$  is the number of sequenced genes.
- $N \times P$  perturbation design matrix  $\mathbf{X}$ , where  $N$  is the number of droplets and  $P$  is the total number of perturbed genes. Here,  $x_{ij}$  represents a binary indicator variable for whether droplet  $i$  contains a guide targeting gene  $j$  (we discuss below how we collapse multiple guides for the same gene).  $\mathbf{X}$  also includes two additional columns corresponding to the presence of a non-targeting control guide and safe-targeting guide, respectively. Cells containing a non-targeting guide are treated as “control” cells (see below), while cells containing a safe-targeting guide are used to test for general effects of genome-targeting guides.

From these data, a  $P \times G$  effect size matrix  $\mathbf{B}$  is estimated, where  $\beta_{ij}$  represents the log fold change of the expression of gene  $j$  relative to control expression when gene  $i$  is perturbed. Two slightly different versions of FR-Perturb were formulated to learn  $\mathbf{B}$  from  $\mathbf{X}$  and  $\mathbf{Y}$  generated from cell and guide pooling, respectively, as follows.

#### *Version 1: Composition in expression space (for cell pooling).*

This scenario arises from cell pooling. The relationship between  $\mathbf{B}$ ,  $\mathbf{X}$ , and  $\mathbf{Y}$  in a given droplet  $i$  is modeled as:

$$E[\mathbf{y}_i] = \frac{1}{g_i} \sum_j x_{ij} \mathbf{c} \exp(\boldsymbol{\beta}_j) \quad (1),$$

where  $\mathbf{y}_i$  is a vector of length  $G$  corresponding to the expression counts of all genes in droplet  $i$ ,  $g_i$  is the number of guides contained in droplet  $i$  (used as a proxy for the number of cells in the droplet),  $x_{ij}$  is a binary scalar indicating whether cell  $i$  contains a guide for gene  $j$ ,  $\mathbf{c}$  is a vector of length  $G$  indicating the expected control expression counts of all genes, and  $\exp(\boldsymbol{\beta}_j)$  is a vector of length  $G$  indicating the fold-change of expression relative to control expression for cells containing a guide for gene  $j$  (with  $\boldsymbol{\beta}_j$  representing the log fold-change). Note that the exp symbol here is used to distinguish fold-changes from log fold-changes, since the latter units are more commonly used to report effect sizes on gene expression. Conceptually, this model reflects the fact that expected expression measured in a droplet containing  $g_i$  cells is the average of the expected expression counts of the individual cells in the droplet (where the latter quantity can be expressed as  $\mathbf{c} \exp(\boldsymbol{\beta}_j)$  for cells containing guide  $j$ ).

In practice, it is advantageous to model the measured expression in each droplet as the *geometric* rather than arithmetic mean of expression of the constituent cells. Simulations with real cells show that the arithmetic versus geometric mean of expression across multiple cells are very similar (**Figure S8A**), but modeling expression counts in a droplet as the latter enables us to perform inference in the space of log fold-changes rather than fold-changes. The former is symmetric around zero (whereas the latter is not) and thus leads to balanced inference of up- versus down-regulation.

Thus, Equation (1) is rewritten as follows:

$$E[\mathbf{y}_i] = \left( \prod_j^P c \exp(\beta_j)^{x_{ij}} \right)^{\frac{1}{g_i}}$$

$$E[\log(\mathbf{y}_i)] = \log(c) + \frac{1}{g_i} \sum_j^P x_{ij} \beta_j \quad (2)$$

Equation (2) can be expressed simply in matrix form as  $E[\mathbf{Y}'] = \mathbf{X}'\mathbf{B}$ , where each row of  $\mathbf{Y}'$ ,  $\mathbf{y}'_i$ , equals  $\log(\mathbf{y}_i) - \log(c)$ , and  $\mathbf{X}'$  is  $\mathbf{X}$  with rows normalized to sum 1. In order to infer  $\mathbf{B}$ ,  $\mathbf{Y}$  is transformed into  $\mathbf{Y}'$  by taking the  $\log(\text{TP10K} + 1)$  of all gene expression counts and subtracting  $\log(c)$  from each row of  $\mathbf{Y}$  (where  $\log(c)$  represents the average  $\log(\text{TP10K} + 1)$  of all genes in cells containing only non-targeting control guides). A pseudocount of 1 is included because the sparse nature of gene expression counts prevents directly taking their logarithm.

Next, the factorize-recover algorithm is applied to  $\mathbf{Y}'$  and  $\mathbf{X}'$  to infer  $\mathbf{B}$ . In the first “factorize” step of factorize-recover, sparse factorization is applied to  $\mathbf{Y}'$  alone using sparse PCA, which produces  $N \times R$  left factor matrix  $\tilde{\mathbf{U}}$  and  $R \times G$  right factor matrix  $\mathbf{W}$ .  $R$  is a hyperparameter that controls the rank of  $\mathbf{Y}'$ . In the second “recover” step, sparse recovery is used to learn  $P \times R$  matrix  $\mathbf{U}$  from the following regression model:  $\tilde{\mathbf{U}} = \mathbf{X}'\mathbf{U}$ , using LASSO applied to each column of  $\tilde{\mathbf{U}}$  (so that one column of  $\mathbf{U}$  is learned at a time). By multiplying  $\mathbf{U}$  by  $\mathbf{W}$  obtained from the factorize step, a  $P \times G$  matrix  $\hat{\mathbf{B}}$  is obtained, which is an estimate of  $\mathbf{B}$ .

In practice, the magnitude of elements of  $\hat{\mathbf{B}}$  was strongly correlated with overall expression level of the downstream gene in control cells. This correlation changed (but was not removed) when varying the arbitrary pseudocount of 1 and/or scale factor of 10,000, suggesting that it was an artifact arising from log-transforming lowly expressed gene expression counts<sup>76</sup>. Indeed, simulations show that the magnitude of effects estimated with FR-Perturb had a negative bias that scaled with the expression level of the downstream gene, with the largest biases observed for the most lowly-expressed genes (**Figure S8C**).

This bias was removed with the following heuristic correction. First, LOESS was used to fit a curve to the plot of effect size magnitude vs. expression level in control cells for all entries of  $\hat{\mathbf{B}}$ . Next, all effect sizes were scaled based on the ratio of their fitted effect size magnitude from LOESS and the fitted effect size magnitude of genes with the highest expression counts ( $\log(\text{average TP10K}) > 2$ ). This procedure removes the global relationship between effect size magnitude and expression level of the downstream gene, while preserving heterogeneity in the

average magnitude of effect sizes on individual downstream genes. In simulations, this procedure produced much less biased effect size estimates than when not scaling (**Figure S8B,C**).

### *Version 2: Composition in log fold change effect size space (for guide pooling)*

For guide pooling data, the relationship between  $\mathbf{B}$ ,  $\mathbf{X}$ , and  $\mathbf{Y}$  in a given droplet  $i$  is modeled as:

$$E[\log(\mathbf{y}_i)] = \log(\mathbf{c}) + \sum_j^P x_{ij} \beta_j \quad (3)$$

The only difference between Equation (2) and Equation (3) is the absence of the normalizing factor  $\frac{1}{g_i}$  in front of the second term of the right side of Equation (3). Inference to learn  $\mathbf{B}$  is performed as in Version 1, with the only difference that the rows of  $\mathbf{X}$  are not normalized to have a sum of 1.

### **Covariates**

Covariates corresponding to the proportion of mitochondrial reads, the total read count per cell, and cell cycle state (as determined by the CellCycleScoring function from the Seurat R package<sup>77</sup>) were accounted for when estimating effect sizes using FR-Perturb, by regressing the covariates out of the expression matrix according to the linear model  $\mathbf{Y}' = \mathbf{CD}$ . Here,  $\mathbf{Y}'$  represents the  $N \times G$  normalized expression matrix (where  $N$  is the number of cells and  $G$  is the number of sequenced genes),  $\mathbf{C}$  represents the  $N \times (C + 1)$  covariate matrix including an intercept term (where  $C$  represents number of covariates with all covariates centered to mean 0), and  $\mathbf{D}$  represents the fitted  $(C + 1) \times G$  matrix of covariate effects on gene expression. All downstream inference was performed on the residual matrix  $\mathbf{Y}_{resid} = \mathbf{Y}' - \mathbf{CD}$ .

### **Hyperparameters for FR-Perturb**

The spams R package<sup>78</sup> was used to perform the steps of factorize-recover, including sparse PCA and LASSO. Three hyperparameters are set in FR-Perturb: the rank  $R$  of  $\mathbf{Y}'$ , a tuning parameter  $\lambda_1$  for sparse PCA during the factorize step (which is the solution of  $\min_{\mathbf{W}} \frac{1}{n} \sum_{i=1}^n \min_{\tilde{\mathbf{u}}_i} \|\mathbf{y}_i - \mathbf{W}\tilde{\mathbf{u}}_i\|_2^2$  so that  $\|\tilde{\mathbf{u}}_i\|_1 \leq \lambda_1$ ), and a tuning parameter  $\lambda_2$  for LASSO during the recover step (which is the solution of  $\min_{\mathbf{u}} \|\tilde{\mathbf{u}} - \mathbf{X}\mathbf{u}\|_2^2$  so that  $\|\mathbf{u}\|_1 \leq \lambda_2$ ). These were set based on maximizing cross-validation  $r^2$  as  $R = 10$ ,  $\lambda_1 = 0.1$ , and  $\lambda_2 = 10$ . Analysis results were not especially sensitive to different values of  $R$ ,  $\lambda_1$ , and  $\lambda_2$  (**Figure S8D-F**).

### **Permutation testing for significance**

Permutation testing was used to obtain two-tailed p-values for elements of  $\hat{\mathbf{B}}$ . To generate an empirical null distribution for each element of  $\hat{\mathbf{B}}$ , samples were permuted (*i.e.*, rows of  $\mathbf{X}$ ) and  $\hat{\mathbf{B}}$  was re-inferred using FR-Perturb for each permutation. Permuting rows of  $\mathbf{X}$  has no impact on the factorize step, since this step does not involve  $\mathbf{X}$  (and the alternative approach of permuting rows of  $\mathbf{Y}$  does not affect the individual factors). Thus, only the recover step was performed and  $\mathbf{U}$  was estimated for each permutation, followed by multiplying the null  $\mathbf{U}$  by  $\mathbf{W}$  obtained from the factorize step to obtain the null  $\hat{\mathbf{B}}$  estimate. In addition, to reduce computational cost, only 500 permutations total were performed. For entries of  $\hat{\mathbf{B}}$  that had p-value = 0 based on these 500 permutations, a skew-t distribution was fit to the empirical null distribution for each entry using the *selm* function from the sn R package, and p-values were then re-computed for these entries

from the fitted distribution. False discovery q-values were computed using the Benjamini-Hochberg procedure applied to the p-values for all entries of  $\hat{\mathbf{B}}$ .

### Inference using negative binomial regression

Using the glmGamPoi R package<sup>79</sup>,  $\mathbf{B}$  was inferred by separately running differential expression analysis for each perturbation (*i.e.*, column of  $\mathbf{X}$ ), where the two groups being compared were droplets containing only non-targeting control guides and droplets containing a guide for the perturbed gene of interest. For droplets containing multiple guides, other guides present in the droplet were ignored when forming these groups. Analytic p-values and false discovery q-values were obtained for all effect sizes from the method output.

### Inference using elastic net

Using the spams R package<sup>78</sup>, the same elastic net inference procedure proposed in Dixit et al.<sup>2</sup> was used to infer  $\mathbf{B}$  from the following models:  $\mathbf{Y}' = \mathbf{X}'\mathbf{B}$  for version 1, and  $\mathbf{Y}' = \mathbf{X}\mathbf{B}$  for version 2 from above with  $\lambda_1 = 0.00025$  and  $\lambda_2 = 0.00025$  (where elastic net finds the solution to  $\min_{\mathbf{y}'} \frac{1}{2} \|\mathbf{y}' - \mathbf{X}\boldsymbol{\beta}\|_2^2 + \lambda_1 \|\boldsymbol{\beta}\|_1 + \frac{\lambda_2}{2} \|\boldsymbol{\beta}\|_2^2$  for each column of  $\mathbf{Y}'$ ), matching the values used in Dixit et al. Other values for the parameters yielded similar results (**Figure S8G**). P-values for all effect sizes were obtained by permuting the rows of  $\mathbf{X}$  a total of 10 times and re-estimating  $\mathbf{B}$  to generate a null distribution across all values of  $\mathbf{B}$ , matching the procedure used in Dixit et al.

### Selecting optimal guide combination for each gene

Four distinct sgRNAs were generated for each perturbed gene. When inferring effect sizes, guides were aggregated by perturbed gene to increase sample size and simplify downstream analyses. When generating the perturbation design matrix  $\mathbf{X}$ , a cell containing any guide for the gene was labelled as receiving a perturbation for the gene. However, sgRNAs have varying efficiency at KO or KD their target gene, and including guides that do not work will add noise to the effect size inference. To retain only sgRNAs that had measurable effects on their target gene, we retained guides with concordant effect size estimates across random sample-wise splits of the data (*i.e.*, the subset of guides to the same gene showing maximal concordance).

Specifically, let  $i$  represent the index of a given perturbed gene, so that  $\mathbf{x}_i$  corresponds to the column of  $\mathbf{X}$  that indicates which cells received perturbation  $i$ , and  $\boldsymbol{\beta}_i$  corresponds to the column of  $\mathbf{B}$  that indicates the effects sizes on all genes' expression from perturbing gene  $i$ . For each  $i$ , 15 different version of  $\mathbf{x}_i$  were generated, corresponding to all possible subsets of the 4 guides. For each version, any cell receiving a guide within the given subset of guides is labelled as containing a perturbation for the gene, while the remaining guides are ignored. Only  $\mathbf{x}_i$  in  $\mathbf{X}$  was modified and the remaining columns were kept the same. Next, the dataset of interest was randomly split in half by samples (cells). FR-Perturb was used to infer effect sizes for all perturbed genes within each half. Then, the  $R^2$  of  $\hat{\boldsymbol{\beta}}_i$  was computed between the two halves (restricting to only effects with an FDR q-value < 0.2), and the specific guide subset that produced that highest  $R^2$  was retained. The same procedure was repeated for each  $i$  to learn the optimal guide combination for each perturbed gene.

### Simulations

Perturb-seq datasets were simulated at various levels of overloading using real expression counts and perturbation effect sizes estimated from our data.

*Simulating cell-pooled data.* To simulate expression data for  $n$  droplets containing  $m$  cells each, the expression of  $n \cdot m$  cells (each containing 1 guide) were first simulated by randomly sampling control cells from our experiment, and scaling their expression counts by the fold change effect sizes of a given perturbed gene (estimated from our conventional knock-out Perturb-seq screen). A 10% probability of receiving a control guide (*i.e.*, no change in expression) was simulated to match the proportion of control guides in the real data. Next, the expression counts of  $m$  cells were randomly averaged at a time to create cell-pooled data.

*Simulating guide-pooled data.* To simulate expression data for  $n$  cells containing  $m$  guides each,  $m$  perturbed genes were randomly selected for each cell, and the expression of a randomly selected control cell was then scaled by the product of the fold change effect sizes of the  $m$  perturbed genes. As before, a 10% probability of receiving a control guide was simulated.

### Clustering and dimensionality reduction

For **Figure 5C**, dimensionality reduction was performed using PCA on the  $\log(\text{TP10K} + 1)$  expression counts of all cells, where the expression values of each gene are scaled and centered to mean 0 and variance 1.

The rows and columns of **Figure 5D** were clustered using Leiden clustering<sup>80</sup>. First, the Euclidian distance between all pairs of genes was calculated by their perturbation effect sizes, and the FindNeighbors function from the Seurat R package<sup>77</sup> was used to compute a shared nearest neighbor graph from these distances ( $k=20$ ), followed by the FindClusters function to perform Leiden clustering on the graph with resolution parameter = 0.5, selected by visual inspection of the resulting clusters. GO enrichment analysis of the genes in the resulting clusters was performed with the ClusterProfiler package<sup>81</sup> with gene sets obtained from the C2 (curated gene sets) and C5 (ontology gene sets) collections of the MSigDB<sup>82</sup>.

### Learning second-order effects for individual perturbation pairs

Second-order interaction effects on gene expression in cell  $i$  with multiple guides were modeled as:

$$E[\log(\mathbf{y}_i)] = \log(\mathbf{c}) + \sum_j^P x_{ij} \boldsymbol{\beta}_j + \sum_j^P \sum_k^P x_{ij} x_{ik} \boldsymbol{\beta}_{jk}$$

Here,  $\log(\mathbf{y}_i)$  is a vector of length  $G$  corresponding to the log expression counts of all genes in droplet  $i$ ,  $x_{ij}$  and  $x_{ik}$  are binary scalars indicating whether cell  $i$  contains a guide for gene  $j$  and/or gene  $k$ ,  $\mathbf{c}$  is a vector of length  $G$  indicating the expected control expression counts of all genes,  $\boldsymbol{\beta}_j$  is a vector of length  $G$  indicating the *first-order* effect size of guide  $j$  on the expression of  $G$  genes, and  $\boldsymbol{\beta}_{jk}$  is a vector of length  $G$  indicating the *second-order* effect size of guides  $j$  and  $k$  on the expression of  $G$  genes. In matrix form, the above can be represented as:

$$E[\mathbf{Y}'] = \mathbf{X}\mathbf{B} + \mathbf{X}_{(2)}\mathbf{B}_{(2)}$$

where each row of  $\mathbf{Y}'$  equals  $\log(\mathbf{y}_i) - \log(\mathbf{c})$ ,  $\mathbf{X}_{(2)}$  is an  $N \times \binom{P}{2}$  indicator matrix for whether each cell contains any of  $\binom{P}{2}$  perturbation pairs, and  $\mathbf{B}_{(2)}$  is a  $\binom{P}{2} \times G$  matrix of second-order

interaction effects.  $\mathbf{B}$  is known from estimating first-order effects previously, which enables the following equation to be written:

$$E[\mathbf{Y}'] = \mathbf{X}_{(2)}\mathbf{B}_{(2)}$$

where  $\mathbf{Y}' = \mathbf{Y} - \mathbf{XB}$ . Finally,  $\mathbf{B}_{(2)}$  is estimated using FR-Perturb in the exact same manner as  $\mathbf{B}$ . To reduce the large size of  $\binom{P}{2}$ , only perturbation pairs that were present in a minimum of 5 cells were included.

When estimating the significance of entries of  $\mathbf{B}_{(2)}$ , the uncertainty in both  $\mathbf{B}$  and  $\mathbf{B}_{(2)}$  must be accounted for, since the latter depends on the former. Thus, when generating a null distribution for the entries of  $\mathbf{B}_{(2)}$ , the rows of both  $\mathbf{X}$  and  $\mathbf{X}_{(2)}$  were permuted and  $\mathbf{B}$  was re-estimated for each permutation.

### Learning second-order effects for perturbation modules

*Intra-modular interactions.* A second-order intra-modular interaction effect was estimated for each co-functional perturbation module  $M$  (i.e., group of perturbed genes) on each co-regulated gene program  $P$  (i.e., group of downstream genes) as follows. For each pair of  $M$  and  $P$ , cells were partitioned into three sets:

- (1) *Control set.* Cells containing only non-targeting control guides or guides for genes without significant effects on  $P$ . The latter group of guides is included to increase sample size, and all these guides are collectively referred to as “control guides”.
- (2) *First-order set.* Cells with exactly one guide in  $M$ , with remaining guides in the cells falling into the “control guide” set.
- (3) *Second-order set.* Cells with exactly two guides in  $M$ , with remaining guides in the cells falling into the “control guide” set.

A mean expression value for  $P$  was computed for each set ( $\mu_0$ ,  $\mu_1$ , and  $\mu_{1,1}$  respectively) as the average standardized  $\log(\text{TP10K} + 1)$  expression of all genes in  $P$  among the cells in the set, with covariates corresponding to read count per cell, percent mitochondrial reads, cell cycle state, and number of guides per cell regressed out of the  $\log(\text{TP10K} + 1)$  expression matrix, and expression standardized to mean 0 and variance 1. The effect size of the first-order set was computed as  $\beta_1 = \mu_1 - \mu_0$  and the interaction effect size of the second-order set as  $\beta_{1,1} = \mu_{1,1} - 2\beta_1 - \mu_0$ . P-values for all interaction effects were computed by permuting the set membership labels of all the cells and recomputing  $\mu_0$ ,  $\beta_1$ , and  $\beta_{1,1}$  for the permuted sets. Standard errors for all interaction effects were computed via bootstrapping, by resampling cells from each of the sets without changing their labels.

*Inter-modular interactions.* Inter-modular interaction effects were computed using a similar approach as above. The 490 total modules were first reduced into 30 disjoint modules using Leiden clustering of a shared nearest neighbor graph defined based on the number of genes shared between gene sets. For two co-functional modules  $M_1$  and  $M_2$ , the first order effects  $\beta_1$  and  $\beta_2$  were computed in the same manner as above. The second-order set was defined as cells with *at least one* guide from each of  $M_1$  and  $M_2$ , with the remaining guides in the cell falling into the “control guide” category, as defined above. The mean expression of the second-order group is  $\mu_{1,2}$ . The interaction effect is defined as  $\beta_{1,2} = \mu_{1,2} - \beta_1 - \beta_2 - \mu_0$  and p-values and standard errors were estimated using permutation testing and bootstrapping, respectively.

## Heritability analyses

Sc-linker<sup>55</sup> was used as previously described to compute a disease heritability enrichment score for each gene set constructed from the KO and KD perturbation effect sizes or perturbation modules and gene programs. Using sc-linker, SNPs were first linked to genes using a combination of histone marks from the Epigenomics Roadmap<sup>83</sup> and the activity-by-contact strategy<sup>84</sup>, then an enrichment score was computed for the SNPs based on the heritability enrichment of the SNPs obtained from stratified LD score regression (S-LDSC<sup>85,86</sup>).

More specifically, for each gene set  $G$ , a set of weights  $A_G = \{a_{G,1}, a_{G,2}, \dots, a_{G,j}\}$  between 0 and 1 was constructed for each SNP based on the confidence of them influencing any gene in  $G$ , following the procedure described in Jagadeesh et al.<sup>55</sup> using activity-by-contact scores<sup>87</sup> and the Epigenomics Roadmap histone marks<sup>83</sup> for whole blood samples. For gene sets defined from membership in perturbation modules (M1-3) or gene programs (P1-4) (**Table S3**), modules/programs were merged between the KO and KD screens. For gene sets defined based on perturbation effects, each gene was weighted by the effect size of the perturbation on the gene, normalized to lie between 0 and 1. A set of weights  $A_{all} = \{a_{all,1}, a_{all,2}, \dots, a_{all,j}\}$  was also constructed, representing the confidence of the SNP influencing any gene across the genome. Next, heritability enrichment estimates  $E_G = \frac{\%h^2(A_G)}{\%SNP(A_G)}$  and  $E_{all} = \frac{\%h^2(A_{all})}{\%SNP(A_{all})}$  were computed for each  $A_G$  and  $A_{all}$ , respectively, using S-LDSC<sup>85,86</sup>. Here,  $\%h^2(A_G) = \frac{\sum_j^M a_{G,j} \beta_j^2}{\sum_j^M \beta_j^2}$  (where  $\beta_j^2$  represents the squared effect size of SNP  $j$  on the phenotype and  $M$  represents the total number of SNPs), and  $\%SNP(A_G) = \frac{\sum_j^M a_{G,j}}{M}$ . Conceptually,  $\%h^2(G)$  represents the fraction of the total genetic effect on the phenotype attributed to SNPs in  $A_G$ , while  $\%SNP(G)$  represents the effective fraction of SNPs that are contained in  $A_G$ . Thus, the ratio  $\frac{\%h^2(G)}{\%SNP(G)}$  is essentially the average effect size magnitude on the phenotype for SNPs in  $A_G$ . Finally, the enrichment score for  $A_G$  was computed as  $E_G - E_{all}$ . Subtracting  $E_{all}$  controls for the baseline level of heritability enrichment for SNPs that influence any gene (since most SNPs do not influence any genes). P-values were obtained for the null hypothesis  $E_G - E_{all} = 0$  using a block jackknife procedure<sup>85</sup>.

## eQTL analyses

Raw genetic data for 432 European individuals and gene expression data for primary monocytes from these individuals profiled 2 hours after treatment with LPS was obtained from Fairfax et al.<sup>26</sup>. For each *cis-trans* gene pair, plink<sup>88</sup> was used to compute marginal association statistics of all SNPs within 1 megabase of the promoter of the *cis* gene with the expression of both the *cis* gene and *trans* gene. All our analyses were restricted to *cis* genes with at least one significant *cis*-eQTL ( $q < 0.05$ ) in the Fairfax dataset. Next, coloc<sup>65</sup> was applied to the association statistics to estimate the posterior probability (with the default prior) that the *cis* and *trans* gene have a shared eQTL within 1 megabase of the *cis* gene, setting a posterior probability threshold of 0.75 to determine significant colocalization (varying this threshold does not change downstream results, **Figure S7D**). The posterior probability that each *cis* gene colocalizes with random *trans* genes was also computed. For all analyses, the top 20 principal components of the gene expression matrix were included as covariates, matching the covariates included by Fairfax et al. in their *trans*-eQTLs analysis and selected based on the fact that they maximize the number of significant *trans*-eQTLs

in Fairfax et al. By restricting the *cis* gene to having a significant eQTL and comparing our effects to random genes while keeping the *cis* gene the same, we control for differences in power for detecting *cis*-by-*trans* eQTLs that arise from differential levels of selective constraint on the *cis* gene. In particular, the *cis* genes selected to be perturbed in our screens include many genes under selective constraint (**Figure S6A**), for which we have decreased power to detect *cis*-by-*trans* eQTLs compared to random *cis* genes.

Bivariate Haseman-Elston regression as implemented in the GCTA software tool<sup>66</sup> was also used to compute the genetic correlation between the expression of the *cis* gene and the *trans* gene when restricting to the region 1 megabase around the promoter of the *cis* gene. Again, the top 20 principal components of the gene expression matrix were included as covariates. The method outputs a genetic correlation estimate  $\hat{r}$  and standard error estimate  $SE(\hat{r})$  for each *cis*-*trans* gene pair. In order to obtain a combined genetic correlation estimate for all downstream genes of a given perturbed gene, all  $\hat{r}$  estimates were first squared and then combined using inverse variance weighing. The variance of  $\hat{r}^2$  was estimated from  $SE(\hat{r})$  using the Delta method:  $Var(\hat{r}^2) \approx 4\hat{r}^2 Var(\hat{r})$ .

### Data and Software Availability

Raw and processed data for all screens were deposited in NCBI's Gene Expression Omnibus under accession number GSE221321. Software implementing FR-Perturb can be found at <https://github.com/douglasyao/FR-Perturb>.

### References

1. Adamson, B., Norman, T.M., Jost, M., Cho, M.Y., Nuñez, J.K., Chen, Y., Villalta, J.E., Gilbert, L.A., Horlbeck, M.A., Hein, M.Y., et al. (2016). A Multiplexed Single-Cell CRISPR Screening Platform Enables Systematic Dissection of the Unfolded Protein Response. *Cell* 167, 1867-1882.e21. 10.1016/j.cell.2016.11.048.
2. Dixit, A., Parnas, O., Li, B., Chen, J., Fulco, C.P., Jerby-Arnon, L., Marjanovic, N.D., Dionne, D., Burks, T., Raychowdhury, R., et al. (2016). Perturb-Seq: Dissecting Molecular Circuits with Scalable Single-Cell RNA Profiling of Pooled Genetic Screens. *Cell* 167, 1853-1866.e17. 10.1016/j.cell.2016.11.038.
3. Jaitin, D.A., Weiner, A., Yofe, I., Lara-Astiaso, D., Keren-Shaul, H., David, E., Salame, T.M., Tanay, A., Oudenaarden, A. van, and Amit, I. (2016). Dissecting Immune Circuits by Linking CRISPR-Pooled Screens with Single-Cell RNA-Seq. *Cell* 167, 1883-1896.e15. 10.1016/j.cell.2016.11.039.
4. Datlinger, P., Rendeiro, A.F., Schmidl, C., Krausgruber, T., Traxler, P., Klughammer, J., Schuster, L.C., Kuchler, A., Alpar, D., and Bock, C. (2017). Pooled CRISPR screening with single-cell transcriptome readout. *Nat. Methods* 14, 297–301. 10.1038/nmeth.4177.
5. Chen, K.H., Boettiger, A.N., Moffitt, J.R., Wang, S., and Zhuang, X. (2015). Spatially resolved, highly multiplexed RNA profiling in single cells. *Science* 348, aaa6090. 10.1126/science.aaa6090.

6. Codeluppi, S., Borm, L.E., Zeisel, A., La Manno, G., van Lunteren, J.A., Svensson, C.I., and Linnarsson, S. (2018). Spatial organization of the somatosensory cortex revealed by osmFISH. *Nat. Methods* 15, 932–935. 10.1038/s41592-018-0175-z.
7. Wang, X., Allen, W.E., Wright, M.A., Sylwestrak, E.L., Samusik, N., Vesuna, S., Evans, K., Liu, C., Ramakrishnan, C., Liu, J., et al. (2018). Three-dimensional intact-tissue sequencing of single-cell transcriptional states. *Science* 361, eaat5691. 10.1126/science.aat5691.
8. Jin, X., Simmons, S.K., Guo, A., Shetty, A.S., Ko, M., Nguyen, L., Jokhi, V., Robinson, E., Oyler, P., Curry, N., et al. (2020). In vivo Perturb-Seq reveals neuronal and glial abnormalities associated with autism risk genes. *Science* 370. 10.1126/science.aaz6063.
9. Fleck, J.S., Jansen, S.M.J., Wollny, D., Zenk, F., Seimiya, M., Jain, A., Okamoto, R., Santel, M., He, Z., Camp, J.G., et al. (2022). Inferring and perturbing cell fate regulomes in human brain organoids. *Nature*, 1–8. 10.1038/s41586-022-05279-8.
10. Paulsen, B., Velasco, S., Kedaigle, A.J., Piloni, M., Quadrato, G., Deo, A.J., Adiconis, X., Uzquiano, A., Sartore, R., Yang, S.M., et al. (2022). Autism genes converge on asynchronous development of shared neuron classes. *Nature* 602, 268–273. 10.1038/s41586-021-04358-6.
11. Replogle, J.M., Saunders, R.A., Pogson, A.N., Hussmann, J.A., Lenail, A., Guna, A., Mascibroda, L., Wagner, E.J., Adelman, K., Lithwick-Yanai, G., et al. (2022). Mapping information-rich genotype-phenotype landscapes with genome-scale Perturb-seq. *Cell* 185, 2559-2575.e28. 10.1016/j.cell.2022.05.013.
12. Freimer, J.W., Shaked, O., Naqvi, S., Sinnott-Armstrong, N., Kathiria, A., Garrido, C.M., Chen, A.F., Cortez, J.T., Greenleaf, W.J., Pritchard, J.K., et al. (2022). Systematic discovery and perturbation of regulatory genes in human T cells reveals the architecture of immune networks. *Nat. Genet.* 54, 1133–1144. 10.1038/s41588-022-01106-y.
13. Frangieh, C.J., Melms, J.C., Thakore, P.I., Geiger-Schuller, K.R., Ho, P., Luoma, A.M., Cleary, B., Jerby-Arnon, L., Malu, S., Cuoco, M.S., et al. (2021). Multimodal pooled Perturb-CITE-seq screens in patient models define mechanisms of cancer immune evasion. *Nat. Genet.* 53, 332–341. 10.1038/s41588-021-00779-1.
14. Norman, T.M., Horlbeck, M.A., Replogle, J.M., Ge, A.Y., Xu, A., Jost, M., Gilbert, L.A., and Weissman, J.S. (2019). Exploring genetic interaction manifolds constructed from rich single-cell phenotypes. *Science* 365, 786–793. 10.1126/science.aax4438.
15. Datlinger, P., Rendeiro, A.F., Boenke, T., Senekowitsch, M., Krausgruber, T., Barreca, D., and Bock, C. (2021). Ultra-high-throughput single-cell RNA sequencing and perturbation screening with combinatorial fluidic indexing. *Nat. Methods* 18, 635–642. 10.1038/s41592-021-01153-z.
16. Gasperini, M., Hill, A.J., McFaline-Figueroa, J.L., Martin, B., Kim, S., Zhang, M.D., Jackson, D., Leith, A., Schreiber, J., Noble, W.S., et al. (2019). A Genome-wide Framework

- for Mapping Gene Regulation via Cellular Genetic Screens. *Cell* 176, 377-390.e19. 10.1016/j.cell.2018.11.029.
17. Candes, E.J., and Wakin, M.B. (2008). An Introduction To Compressive Sampling. *IEEE Signal Process. Mag.* 25, 21–30. 10.1109/MSP.2007.914731.
18. Candes, E.J., Romberg, J., and Tao, T. (2006). Robust uncertainty principles: exact signal reconstruction from highly incomplete frequency information. *IEEE Trans. Inf. Theory* 52, 489–509. 10.1109/TIT.2005.862083.
19. Donoho, D.L. (2006). Compressed sensing. *IEEE Trans. Inf. Theory* 52, 1289–1306. 10.1109/TIT.2006.871582.
20. Cleary, B., Cong, L., Cheung, A., Lander, E.S., and Regev, A. (2017). Efficient Generation of Transcriptomic Profiles by Random Composite Measurements. *Cell* 171, 1424-1436.e18. 10.1016/j.cell.2017.10.023.
21. Cleary, B., Simonton, B., Bezney, J., Murray, E., Alam, S., Sinha, A., Habibi, E., Marshall, J., Lander, E.S., Chen, F., et al. (2021). Compressed sensing for highly efficient imaging transcriptomics. *Nat. Biotechnol.*, 1–7. 10.1038/s41587-021-00883-x.
22. Sharan, V., Tai, K.S., Bailis, P., and Valiant, G. (2019). Compressed Factorization: Fast and Accurate Low-Rank Factorization of Compressively-Sensed Data. In *Proceedings of the 36th International Conference on Machine Learning (PMLR)*, pp. 5690–5700.
23. Yeung, K.Y., and Ruzzo, W.L. (2001). Principal component analysis for clustering gene expression data. *Bioinformatics* 17, 763–774. 10.1093/bioinformatics/17.9.763.
24. Brunet, J.-P., Tamayo, P., Golub, T.R., and Mesirov, J.P. (2004). Metagenes and molecular pattern discovery using matrix factorization. *Proc. Natl. Acad. Sci.* 101, 4164–4169.
25. Parnas, O., Jovanovic, M., Eisenhaure, T.M., Herbst, R.H., Dixit, A., Ye, C.J., Przybylski, D., Platt, R.J., Tirosh, I., Sanjana, N.E., et al. (2015). A Genome-wide CRISPR Screen in Primary Immune Cells to Dissect Regulatory Networks. *Cell* 162, 675–686. 10.1016/j.cell.2015.06.059.
26. Fairfax, B.P., Humburg, P., Makino, S., Naranbhai, V., Wong, D., Lau, E., Jostins, L., Plant, K., Andrews, R., McGee, C., et al. (2014). Innate Immune Activity Conditions the Effect of Regulatory Variants upon Monocyte Gene Expression. *Science* 343, 1246949. 10.1126/science.1246949.
27. Chanput, W., Mes, J.J., and Wichers, H.J. (2014). THP-1 cell line: An in vitro cell model for immune modulation approach. *Int. Immunopharmacol.* 23, 37–45. 10.1016/j.intimp.2014.08.002.
28. Aguirre, A.J., Meyers, R.M., Weir, B.A., Vazquez, F., Zhang, C.-Z., Ben-David, U., Cook, A., Ha, G., Harrington, W.F., Doshi, M.B., et al. (2016). Genomic Copy Number

- Dictates a Gene-Independent Cell Response to CRISPR/Cas9 Targeting. *Cancer Discov.* 6, 914–929. 10.1158/2159-8290.CD-16-0154.
29. Srivatsan, S.R., McFaline-Figueroa, J.L., Ramani, V., Saunders, L., Cao, J., Packer, J., Pliner, H.A., Jackson, D.L., Daza, R.M., Christiansen, L., et al. (2020). Massively multiplex chemical transcriptomics at single-cell resolution. *Science* 367, 45–51. 10.1126/science.aax6234.
30. Zetsche, B., Gootenberg, J.S., Abudayyeh, O.O., Slaymaker, I.M., Makarova, K.S., Essletzbichler, P., Volz, S.E., Joung, J., van der Oost, J., Regev, A., et al. (2015). Cpf1 Is a Single RNA-Guided Endonuclease of a Class 2 CRISPR-Cas System. *Cell* 163, 759–771. 10.1016/j.cell.2015.09.038.
31. Campa, C.C., Weisbach, N.R., Santinha, A.J., Incarnato, D., and Platt, R.J. (2019). Multiplexed genome engineering by Cas12a and CRISPR arrays encoded on single transcripts. *Nat. Methods* 16, 887–893. 10.1038/s41592-019-0508-6.
32. Rosenbluh, J., Xu, H., Harrington, W., Gill, S., Wang, X., Vazquez, F., Root, D.E., Tsherniak, A., and Hahn, W.C. (2017). Complementary information derived from CRISPR Cas9 mediated gene deletion and suppression. *Nat. Commun.* 8, 15403. 10.1038/ncomms15403.
33. Brubaker, S.W., Bonham, K.S., Zanoni, I., and Kagan, J.C. (2015). Innate Immune Pattern Recognition: A Cell Biological Perspective. *Annu. Rev. Immunol.* 33, 257–290. 10.1146/annurev-immunol-032414-112240.
34. Palucka, A.K., Blanck, J.-P., Bennett, L., Pascual, V., and Banchereau, J. (2005). Cross-regulation of TNF and IFN- $\alpha$  in autoimmune diseases. *Proc. Natl. Acad. Sci.* 102, 3372–3377. 10.1073/pnas.0408506102.
35. Mavragani, C.P., Niewold, T.B., Moutsopoulos, N.M., Pillemer, S.R., Wahl, S.M., and Crow, M.K. (2007). Augmented interferon-alpha pathway activation in patients with Sjögren's syndrome treated with etanercept. *Arthritis Rheum.* 56, 3995–4004. 10.1002/art.23062.
36. Dorrington, M.G., and Fraser, I.D.C. (2019). NF- $\kappa$ B Signaling in Macrophages: Dynamics, Crosstalk, and Signal Integration. *Front. Immunol.* 10.
37. Wang, N., Liang, H., and Zen, K. (2014). Molecular Mechanisms That Influence the Macrophage M1–M2 Polarization Balance. *Front. Immunol.* 5.
38. Komura, T., Sakai, Y., Honda, M., Takamura, T., Wada, T., and Kaneko, S. (2013). ER stress induced impaired TLR signaling and macrophage differentiation of human monocytes. *Cell. Immunol.* 282, 44–52. 10.1016/j.cellimm.2013.04.006.
39. Plataniias, L.C. (2005). Mechanisms of type-I- and type-II-interferon-mediated signalling. *Nat. Rev. Immunol.* 5, 375–386. 10.1038/nri1604.

40. Carballo, E., Lai, W.S., and Blakeshear, P.J. (1998). Feedback Inhibition of Macrophage Tumor Necrosis Factor- $\alpha$  Production by Tristetraprolin. *Science* 281, 1001–1005. 10.1126/science.281.5379.1001.
41. Trompouki, E., Hatzivassiliou, E., Tschritzis, T., Farmer, H., Ashworth, A., and Mosialos, G. (2003). CYLD is a deubiquitinating enzyme that negatively regulates NF- $\kappa$ B activation by TNFR family members. *Nature* 424, 793–796. 10.1038/nature01803.
42. Shembade, N., Ma, A., and Harhaj, E.W. (2010). Inhibition of NF- $\kappa$ B Signaling by A20 Through Disruption of Ubiquitin Enzyme Complexes. *Science* 327, 1135–1139. 10.1126/science.1182364.
43. Wertz, I.E., O'Rourke, K.M., Zhang, Z., Dornan, D., Arnott, D., Deshaies, R.J., and Dixit, V.M. (2004). Human De-Etiolated-1 Regulates c-Jun by Assembling a CUL4A Ubiquitin Ligase. *Science* 303, 1371–1374. 10.1126/science.1093549.
44. Kiss-Toth, E., Bagstaff, S.M., Sung, H.Y., Jozsa, V., Dempsey, C., Caunt, J.C., Oxley, K.M., Wyllie, D.H., Polgar, T., Harte, M., et al. (2004). Human Tribbles, a Protein Family Controlling Mitogen-activated Protein Kinase Cascades \*. *J. Biol. Chem.* 279, 42703–42708. 10.1074/jbc.M407732200.
45. Scholz-Starke, J., and Cesca, F. (2016). Stepping Out of the Shade: Control of Neuronal Activity by the Scaffold Protein Kidins220/ARMS. *Front. Cell. Neurosci.* 10.
46. Huttlin, E.L., Bruckner, R.J., Navarrete-Perea, J., Cannon, J.R., Baltier, K., Gebreab, F., Gygi, M.P., Thornock, A., Zarraga, G., Tam, S., et al. (2021). Dual proteome-scale networks reveal cell-specific remodeling of the human interactome. *Cell* 184, 3022–3040.e28. 10.1016/j.cell.2021.04.011.
47. Amiot, F., Boussadia, O., Cases, S., Fitting, C., Lebastard, M., Cavaillon, J.-M., Milon, G., and Dautry, F. (1997). Mice heterozygous for a deletion of the tumor necrosis factor- $\alpha$  and lymphotoxin- $\alpha$  genes: biological importance of a nonlinear response of tumor necrosis factor- $\alpha$  to gene dosage. *Eur. J. Immunol.* 27, 1035–1042. 10.1002/eji.1830270434.
48. Simon, A., Park, H., Maddipati, R., Lobito, A.A., Bulua, A.C., Jackson, A.J., Chae, J.J., Ettinger, R., de Koning, H.D., Cruz, A.C., et al. (2010). Concerted action of wild-type and mutant TNF receptors enhances inflammation in TNF receptor 1-associated periodic fever syndrome. *Proc. Natl. Acad. Sci.* 107, 9801–9806. 10.1073/pnas.0914118107.
49. Segrè, D., DeLuna, A., Church, G.M., and Kishony, R. (2005). Modular epistasis in yeast metabolism. *Nat. Genet.* 37, 77–83. 10.1038/ng1489.
50. Costanzo, M., Baryshnikova, A., Bellay, J., Kim, Y., Spear, E.D., Sevier, C.S., Ding, H., Koh, J.L.Y., Toufighi, K., Mostafavi, S., et al. (2010). The Genetic Landscape of a Cell. *Science* 327, 425–431. 10.1126/science.1180823.

51. Lang, K.S., Burow, A., Kurrer, M., Lang, P.A., and Recher, M. (2007). The role of the innate immune response in autoimmune disease. *J. Autoimmun.* *29*, 206–212. 10.1016/j.jaut.2007.07.018.
52. O'Connor, L.J., Schoech, A.P., Hormozdiari, F., Gazal, S., Patterson, N., and Price, A.L. (2019). Extreme Polygenicity of Complex Traits Is Explained by Negative Selection. *Am. J. Hum. Genet.* *105*, 456–476. 10.1016/j.ajhg.2019.07.003.
53. Lek, M., Karczewski, K.J., Minikel, E.V., Samocha, K.E., Banks, E., Fennell, T., O'Donnell-Luria, A.H., Ware, J.S., Hill, A.J., Cummings, B.B., et al. (2016). Analysis of protein-coding genetic variation in 60,706 humans. *Nature* *536*, 285–291. 10.1038/nature19057.
54. Karczewski, K.J., Francioli, L.C., Tiao, G., Cummings, B.B., Alföldi, J., Wang, Q., Collins, R.L., Laricchia, K.M., Ganna, A., Birnbaum, D.P., et al. (2020). The mutational constraint spectrum quantified from variation in 141,456 humans. *Nature* *581*, 434–443. 10.1038/s41586-020-2308-7.
55. Jagadeesh, K.A., Dey, K.K., Montoro, D.T., Mohan, R., Gazal, S., Engreitz, J.M., Xavier, R.J., Price, A.L., and Regev, A. (2022). Identifying disease-critical cell types and cellular processes by integrating single-cell RNA-sequencing and human genetics. *Nat. Genet.* *54*, 1479–1492. 10.1038/s41588-022-01187-9.
56. Morris, J.A., Daniloski, Z., Domingo, J., Barry, T., Ziosi, M., Glinos, D.A., Hao, S., Mimitou, E.P., Smibert, P., Roeder, K., et al. (2021). Discovery of target genes and pathways of blood trait loci using pooled CRISPR screens and single cell RNA sequencing. 2021.04.07.438882. 10.1101/2021.04.07.438882.
57. Graustein, A., Misch, E.A., Musvosvi, M., Shey, M., Shah, J., Wells, R., Hanekom, W., Hatherill, M., Scriba, T., and Hawn, T. (2016). HSP90B1 Regulates TLR-dependent Monocyte Signaling and its Common Variants are Associated with BCG-specific T-cell Responses and Protection from Pediatric TB Disease. *J. Immunol.* *196*, 200.18-200.18.
58. Casey, S.C., Tong, L., Li, Y., Do, R., Walz, S., Fitzgerald, K.N., Gouw, A.M., Baylot, V., Gütgemann, I., Eilers, M., et al. (2016). MYC regulates the antitumor immune response through CD47 and PD-L1. *Science* *352*, 227–231. 10.1126/science.aac9935.
59. Kortlever, R.M., Sodir, N.M., Wilson, C.H., Burkhart, D.L., Pellegrinet, L., Swigart, L.B., Littlewood, T.D., and Evan, G.I. (2017). Myc Cooperates with Ras by Programming Inflammation and Immune Suppression. *Cell* *171*, 1301-1315.e14. 10.1016/j.cell.2017.11.013.
60. Garcia-Etxebarria, K., Bracho, M.A., Galán, J.C., Pumarola, T., Castilla, J., Lejarazu, R.O. de, Rodríguez-Dominguez, M., Quintela, I., Bonet, N., Garcia-Garcera, M., et al. (2015). No Major Host Genetic Risk Factor Contributed to A(H1N1)2009 Influenza Severity. *PLOS ONE* *10*, e0135983. 10.1371/journal.pone.0135983.
61. Han, H., Cho, J.-W., Lee, S., Yun, A., Kim, H., Bae, D., Yang, S., Kim, C.Y., Lee, M., Kim, E., et al. (2018). TRRUST v2: an expanded reference database of human and mouse

- transcriptional regulatory interactions. *Nucleic Acids Res.* 46, D380–D386.  
10.1093/nar/gkx1013.
62. Liu, X., Li, Y.I., and Pritchard, J.K. (2019). Trans Effects on Gene Expression Can Drive Omnigenic Inheritance. *Cell* 177, 1022–1034.e6. 10.1016/j.cell.2019.04.014.
63. Võsa, U., Claringbould, A., Westra, H.-J., Bonder, M.J., Deelen, P., Zeng, B., Kirsten, H., Saha, A., Kreuzhuber, R., Yazar, S., et al. (2021). Large-scale cis- and trans-eQTL analyses identify thousands of genetic loci and polygenic scores that regulate blood gene expression. *Nat. Genet.* 53, 1300–1310. 10.1038/s41588-021-00913-z.
64. Westra, H.-J., Peters, M.J., Esko, T., Yaghootkar, H., Schurmann, C., Kettunen, J., Christiansen, M.W., Fairfax, B.P., Schramm, K., Powell, J.E., et al. (2013). Systematic identification of trans eQTLs as putative drivers of known disease associations. *Nat. Genet.* 45, 1238–1243. 10.1038/ng.2756.
65. Giambartolomei, C., Vukcevic, D., Schadt, E.E., Franke, L., Hingorani, A.D., Wallace, C., and Plagnol, V. (2014). Bayesian Test for Colocalisation between Pairs of Genetic Association Studies Using Summary Statistics. *PLOS Genet.* 10, e1004383. 10.1371/journal.pgen.1004383.
66. Yang, J., Lee, S.H., Goddard, M.E., and Visscher, P.M. (2011). GCTA: A Tool for Genome-wide Complex Trait Analysis. *Am. J. Hum. Genet.* 88, 76–82. 10.1016/j.ajhg.2010.11.011.
67. Lukowski, S.W., Lloyd-Jones, L.R., Holloway, A., Kirsten, H., Hemani, G., Yang, J., Small, K., Zhao, J., Metspalu, A., Dermitzakis, E.T., et al. (2017). Genetic correlations reveal the shared genetic architecture of transcription in human peripheral blood. *Nat. Commun.* 8, 483. 10.1038/s41467-017-00473-z.
68. Umans, B.D., Battle, A., and Gilad, Y. (2021). Where Are the Disease-Associated eQTLs? *Trends Genet.* 37, 109–124. 10.1016/j.tig.2020.08.009.
69. Simmons, S.K., Lithwick-Yanai, G., Adiconis, X., Oberstrass, F., Iremadze, N., Geiger-Schuller, K., Thakore, P.I., Frangieh, C.J., Barad, O., Almogy, G., et al. (2022). Mostly natural sequencing-by-synthesis for scRNA-seq using Ultima sequencing. *Nat. Biotechnol.*, 1–8. 10.1038/s41587-022-01452-6.
70. Schraivogel, D., Gschwind, A.R., Milbank, J.H., Leonce, D.R., Jakob, P., Mathur, L., Korbel, J.O., Merten, C.A., Velten, L., and Steinmetz, L.M. (2020). Targeted Perturb-seq enables genome-scale genetic screens in single cells. *Nat. Methods* 17, 629–635. 10.1038/s41592-020-0837-5.
71. O'Connor, L.J. (2021). The distribution of common-variant effect sizes. *Nat. Genet.* 53, 1243–1249. 10.1038/s41588-021-00901-3.
72. Yazar, S., Alquicira-Hernandez, J., Wing, K., Senabouth, A., Gordon, M.G., Andersen, S., Lu, Q., Rowson, A., Taylor, T.R.P., Clarke, L., et al. (2022). Single-cell eQTL mapping

- identifies cell type-specific genetic control of autoimmune disease. *Science*. 10.1126/science.abf3041.
73. Huang, H., Fang, M., Jostins, L., Umićević Mirkov, M., Boucher, G., Anderson, C.A., Andersen, V., Cleyneen, I., Cortes, A., Crins, F., et al. (2017). Fine-mapping inflammatory bowel disease loci to single-variant resolution. *Nature* 547, 173–178. 10.1038/nature22969.
74. Amberger, J.S., Bocchini, C.A., Scott, A.F., and Hamosh, A. (2019). OMIM.org: leveraging knowledge across phenotype–gene relationships. *Nucleic Acids Res.* 47, D1038–D1043. 10.1093/nar/gky1151.
75. Buniello, A., MacArthur, J.A.L., Cerezo, M., Harris, L.W., Hayhurst, J., Malangone, C., McMahon, A., Morales, J., Mountjoy, E., Sollis, E., et al. (2019). The NHGRI-EBI GWAS Catalog of published genome-wide association studies, targeted arrays and summary statistics 2019. *Nucleic Acids Res.* 47, D1005–D1012. 10.1093/nar/gky1120.
76. O’Hara, R., and Kotze, J. (2010). Do not log-transform count data. *Nat. Preced.*, 1–1. 10.1038/npre.2010.4136.1.
77. Stuart, T., Butler, A., Hoffman, P., Hafemeister, C., Papalexi, E., Mauck, W.M., Hao, Y., Stoeckius, M., Smibert, P., and Satija, R. (2019). Comprehensive Integration of Single-Cell Data. *Cell* 177, 1888–1902.e21. 10.1016/j.cell.2019.05.031.
78. Mairal, J., Bach, F., Ponce, J., and Sapiro, G. (2010). Online Learning for Matrix Factorization and Sparse Coding. *J. Mach. Learn. Res.* 11, 19–60.
79. Ahlmann-Eltze, C., and Huber, W. (2020). glmGamPoi: fitting Gamma-Poisson generalized linear models on single cell count data. *Bioinformatics* 36, 5701–5702. 10.1093/bioinformatics/btaa1009.
80. Traag, V.A., Waltman, L., and van Eck, N.J. (2019). From Louvain to Leiden: guaranteeing well-connected communities. *Sci. Rep.* 9, 5233. 10.1038/s41598-019-41695-z.
81. Yu, G., Wang, L.-G., Han, Y., and He, Q.-Y. (2012). clusterProfiler: an R Package for Comparing Biological Themes Among Gene Clusters. *OMICS J. Integr. Biol.* 16, 284–287. 10.1089/omi.2011.0118.
82. Liberzon, A., Subramanian, A., Pinchback, R., Thorvaldsdóttir, H., Tamayo, P., and Mesirov, J.P. (2011). Molecular signatures database (MSigDB) 3.0. *Bioinformatics* 27, 1739–1740. 10.1093/bioinformatics/btr260.
83. Kundaje, A., Meuleman, W., Ernst, J., Bilenky, M., Yen, A., Heravi-Moussavi, A., Kheradpour, P., Zhang, Z., Wang, J., Ziller, M.J., et al. (2015). Integrative analysis of 111 reference human epigenomes. *Nature* 518, 317–330. 10.1038/nature14248.
84. Fulco, C.P., Nasser, J., Jones, T.R., Munson, G., Bergman, D.T., Subramanian, V., Grossman, S.R., Anyoha, R., Doughty, B.R., Patwardhan, T.A., et al. (2019). Activity-by-

- contact model of enhancer–promoter regulation from thousands of CRISPR perturbations. *Nat. Genet.* *51*, 1664–1669. 10.1038/s41588-019-0538-0.
85. Finucane, H.K., Bulik-Sullivan, B., Gusev, A., Trynka, G., Reshef, Y., Loh, P.-R., Anttila, V., Xu, H., Zang, C., Farh, K., et al. (2015). Partitioning heritability by functional annotation using genome-wide association summary statistics. *Nat. Genet.* *47*, 1228–1235. 10.1038/ng.3404.
86. Gazal, S., Finucane, H.K., Furlotte, N.A., Loh, P.-R., Palamara, P.F., Liu, X., Schoech, A., Bulik-Sullivan, B., Neale, B.M., Gusev, A., et al. (2017). Linkage disequilibrium–dependent architecture of human complex traits shows action of negative selection. *Nat. Genet.* *49*, 1421–1427. 10.1038/ng.3954.
87. Nasser, J., Bergman, D.T., Fulco, C.P., Guckelberger, P., Doughty, B.R., Patwardhan, T.A., Jones, T.R., Nguyen, T.H., Ulirsch, J.C., Lekschas, F., et al. (2021). Genome-wide enhancer maps link risk variants to disease genes. *Nature* *593*, 238–243. 10.1038/s41586-021-03446-x.
88. Purcell, S., Neale, B., Todd-Brown, K., Thomas, L., Ferreira, M.A.R., Bender, D., Maller, J., Sklar, P., de Bakker, P.I.W., Daly, M.J., et al. (2007). PLINK: A Tool Set for Whole-Genome Association and Population-Based Linkage Analyses. *Am. J. Hum. Genet.* *81*, 559–575.