

Predicting cell population-specific gene expression from genomic sequence

Michielsen, Lieke; Reinders, Marcel J. T.; Mahfouz, Ahmed

DOI

[10.3389/fbinf.2024.1347276](https://doi.org/10.3389/fbinf.2024.1347276)

Publication date

2024

Document Version

Final published version

Published in

Frontiers in Bioinformatics

Citation (APA)

Michielsen, L., Reinders, M. J. T., & Mahfouz, A. (2024). Predicting cell population-specific gene expression from genomic sequence. *Frontiers in Bioinformatics*, 4, Article 1347276.
<https://doi.org/10.3389/fbinf.2024.1347276>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.



OPEN ACCESS

EDITED BY

Nagarajan Raju,
Emory University, United States

REVIEWED BY

Disha Sharma,
Stanford Healthcare, United States
Stefan Semrau,
New York Stem Cell Foundation, United States

*CORRESPONDENCE

Ahmed Mahfouz,
✉ a.mahfouz@lumc.nl

RECEIVED 30 November 2023

ACCEPTED 23 January 2024

PUBLISHED 04 March 2024

CITATION

Michielsen L, Reinders MJT and Mahfouz A
(2024), Predicting cell population-specific gene
expression from genomic sequence.
Front. Bioinform. 4:1347276.
doi: 10.3389/fbinf.2024.1347276

COPYRIGHT

© 2024 Michielsen, Reinders and Mahfouz. This
is an open-access article distributed under the
terms of the [Creative Commons Attribution
License \(CC BY\)](#). The use, distribution or
reproduction in other forums is permitted,
provided the original author(s) and the
copyright owner(s) are credited and that the
original publication in this journal is cited, in
accordance with accepted academic practice.
No use, distribution or reproduction is
permitted which does not comply with these
terms.

Predicting cell population-specific gene expression from genomic sequence

Lieke Michielsen^{1,2,3}, Marcel J. T. Reinders^{1,2,3} and
Ahmed Mahfouz^{1,2,3*}

¹Department of Human Genetics, Leiden University Medical Center, Leiden, Netherlands, ²Leiden
Computational Biology Center, Leiden University Medical Center, Leiden, Netherlands, ³Delft
Bioinformatics Lab, Delft University of Technology, Delft, Netherlands

Most regulatory elements, especially enhancer sequences, are cell population-specific. One could even argue that a distinct set of regulatory elements is what defines a cell population. However, discovering which non-coding regions of the DNA are essential in which context, and as a result, which genes are expressed, is a difficult task. Some computational models tackle this problem by predicting gene expression directly from the genomic sequence. These models are currently limited to predicting bulk measurements and mainly make tissue-specific predictions. Here, we present a model that leverages single-cell RNA-sequencing data to predict gene expression. We show that cell population-specific models outperform tissue-specific models, especially when the expression profile of a cell population and the corresponding tissue are dissimilar. Further, we show that our model can prioritize GWAS variants and learn motifs of transcription factor binding sites. We envision that our model can be useful for delineating cell population-specific regulatory elements.

KEYWORDS

sequence to prediction models, single-cell RNA-sequencing, gene expression prediction, transcriptional regulation, cell populations

1 Introduction

In multicellular organisms, every cell has the same DNA apart from somatic mutations. Yet its function and the related proteins and genes expressed vary enormously. This is among others caused by transcriptional and epigenetic regulation. Proteins that bind the DNA sequence around the transcription start site (TSS) control whether a gene is transcribed in a cell (Vaquerizas et al., 2009; Lambert et al., 2018). Which transcription factors, and thus which DNA binding motifs, are essential differ per cell population (Vaquerizas et al., 2009; Lambert et al., 2018; Nott et al., 2019; Janssens et al., 2022). As such, mutations in regulatory regions might affect specific tissues or cell populations differently. Improving our understanding of these regulatory mechanisms will help us relate genomic functions to a phenotype.

For example, while promoter sequences are identical across the four major human brain cell populations (neurons, oligodendrocytes, astrocytes, and microglia), almost all enhancer sequences, the regions in the DNA where a transcription factor binds, are cell population-specific (Nott et al., 2019). These population-specific regulatory elements are discovered by combining single-cell measurements of different data types, including chromatin

accessibility, ChIP-seq, and DNA methylation. Bakken et al. (2021), for instance, identified differentially methylated and differentially accessible regions across neuronal cell populations in the human brain, albeit with little overlap. This emphasizes the complexity of transcriptional regulation and the need for more measurements to fully resolve these mechanisms at the cell population-specific level.

An alternative approach would be to train a computational model that directly predicts gene expression from the genomic sequence around the TSS. This way, we can learn which regulatory elements are important for transcriptional regulation in different contexts. Several computational methods have been developed for this task (Kelley et al., 2016; Kelley et al., 2018; Zhou et al., 2018; Agarwal and Shendure, 2020; Zhang et al., 2020; Avsec et al., 2021a; Avsec et al., 2021b). These methods have in common that they one-hot encode the DNA sequence and input this to either a convolutional neural network (CNN) or transformer. ExPecto, Xpresso, and ExpResNet predict expression measurements from bulk RNA-sequencing, while Basset, Basenji, BPNet, and the Enformer model predict regulatory signals, such as cap analysis gene expression (CAGE) reads or TF binding from CHIP-nexus.

A promising application of these models is to prioritize variants that have been identified using genome-wide association studies (GWAS) (Kelley et al., 2016; Wesolowska-Andersen et al., 2020). Using GWAS many potential disease-associating variants have been identified (Schizophrenia Working Group of the Psychiatric Genomics Consortium, 2014; Wightman et al., 2021; Yao et al., 2021). Within each locus, however, it is often challenging to pinpoint which variant is causal and which gene is affected by the variant.

These current computational gene prediction models, however, are designed for predicting bulk gene expression data. This means that they are either tissue-specific or could be applied to FACS-sorted cells (Wesolowska-Andersen et al., 2020). Since transcriptional regulation is even more context-specific, the resolution of current methods is not sufficient for heterogeneous tissues where single-cell RNA-sequencing (scRNA-seq) has revealed hundreds of cell populations (Tasic et al., 2016; Tasic et al., 2018; Bakken et al., 2021). To increase the resolution, the models would ideally be trained on scRNA-seq data.

Here, we present scXpresso, a deep learning model that uses a CNN to learn cell population-specific expression in scRNA-seq data from genomic sequences. Since single-cell and bulk data have different characteristics and distributions, we explored whether this type of model is suitable for single-cell data. We show that 1) cell population-specific models outperform tissue-specific models on several tissues from the Tabula Muris, 2) increasing the resolution improves the predictions for human brain cell populations, and 3) *in silico* saturation mutagenesis of the input sequence can be used to prioritize GWAS variants.

2 Materials and methods

2.1 Architecture of scXpresso

scXpresso is a one-dimensional convolutional neural network (CNN) adapted from the (bulk gene expression-based) Xpresso

model (Agarwal and Shendure, 2020) (Figure 1A; Supplementary Figure S1). The input to the CNN is four channels with the one-hot encoded sequence around the transcription start site (TSS) (7 kb upstream and 3.5 kb downstream). Every channel represents one of the four nucleotides (A, C, T, G). For some positions, the exact nucleotide is not known [e.g., any nucleic acid (N) or a purine nucleotide (R)]. The exact coding scheme for such positions is shown in Supplementary Table S1. The CNN consists of two convolutional layers. The output of the convolutional layers is flattened and concatenated with the half-life time features. Together, this is subsequently fed into a fully connected (FC) layer(s). The output of the FC layers is the aggregated expression per tissue or for each cell population.

Comparing scXpresso to Xpresso, there are three main differences: 1) we designed scXpresso as a multitask model so that it predicts the expression of multiple cell populations simultaneously. 2) We decreased the number of half-life time features from eight to five; the three features we removed (5' UTR, ORF, and 3' UTR GC content) correlated less with half-life time, so we removed them to make the model less complex (Sharova et al., 2009; Spies et al., 2013; Agarwal and Shendure, 2020). Furthermore, removing these three half-life time features from the original Xpresso model did not lower its performance (Supplementary Table S2). 3) For the multitask model, there is only one FC layer. For the other models, which we use to make tissue-specific predictions as a comparison, we used two FC layers.

2.2 Training scXpresso

We split the genes into a train, validation, and test dataset and evaluated using 20-fold cross-validation. These sets are the same across all experiments (i.e., one train, validation, and test set for mouse genes and one for human genes) such that the results of different models can be compared. We update the weights of scXpresso using the Adam optimizer based on the mean square error loss on the training set. The initial learning rate is set to 0.0005 and if the loss on the validation set is not improved from 5 epochs, the learning rate is reduced by a factor of 10. We train the model for 40 epochs and the model with the lowest loss on the validation set is used for evaluation on the test dataset. Since there is always some stochasticity when training a CNN, we always train 5 models and average the predictions. We used the following software packages for training the model: Pytorch (version 1.9.0) (Paszke et al., 2019), CUDA (version 11.1), cuDNN (version 8.0.5.39), and Python (version 3.6.8).

2.3 Datasets

2.3.1 Tabula Muris

The single-cell Tabula Muris data (Schaum et al., 2018) for the five different tissues (gland, spleen, lung, limb muscle, and bone marrow) and two different protocols (10X and FACS-based Smart-seq2) were downloaded from: https://figshare.com/projects/Tabula_Muris_Transcriptomic_characterization_of_20_organisms_and_tissues_

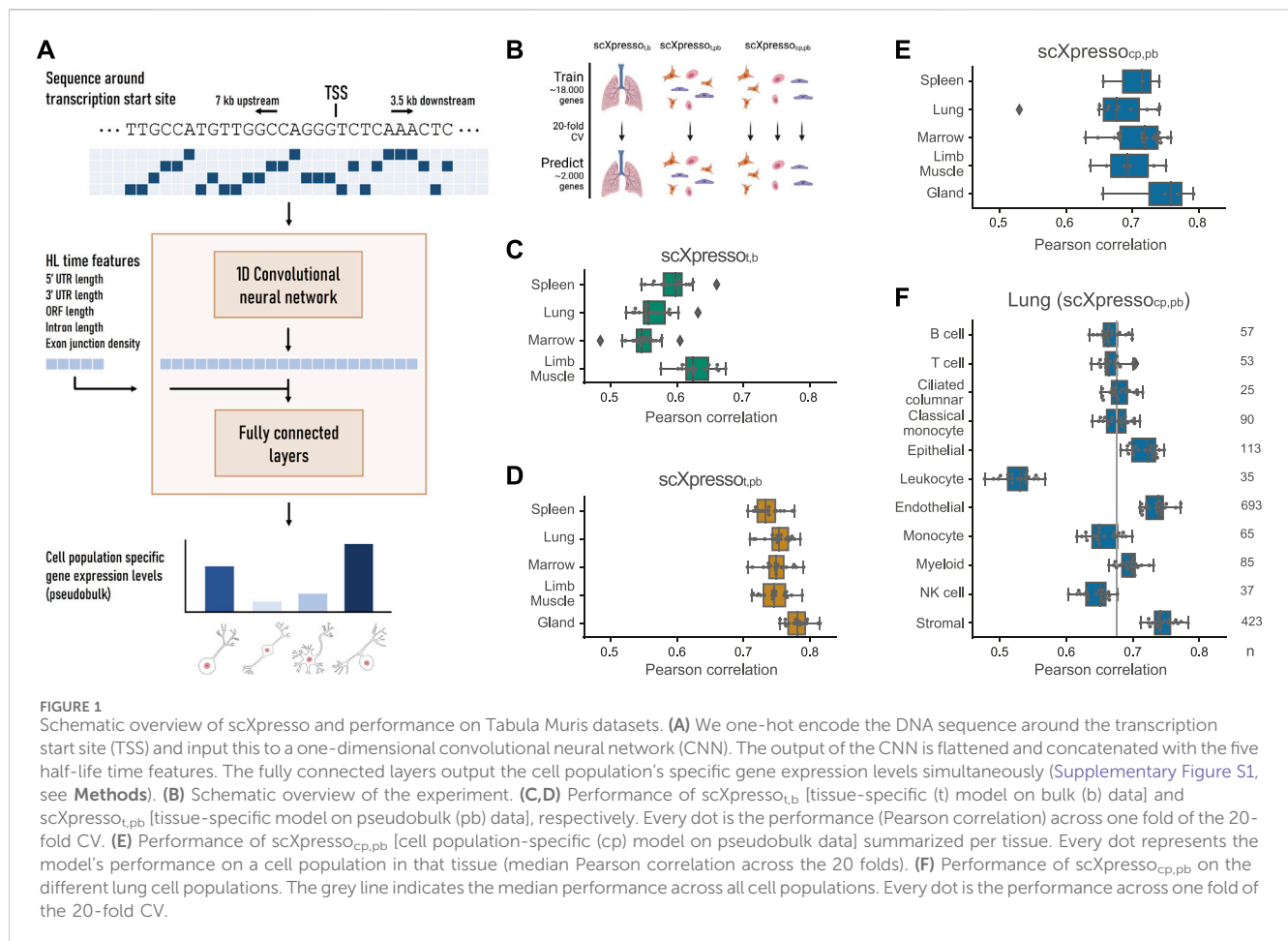


FIGURE 1

Schematic overview of scXpresso and performance on Tabula Muris datasets. (A) We one-hot encode the DNA sequence around the transcription start site (TSS) and input this to a one-dimensional convolutional neural network (CNN). The output of the CNN is flattened and concatenated with the five half-life time features. The fully connected layers output the cell population's specific gene expression levels simultaneously (Supplementary Figure S1, see **Methods**). (B) Schematic overview of the experiment. (C,D) Performance of scXpresso_{t,b} [tissue-specific (t) model on bulk (b) data] and scXpresso_{t,pb} [tissue-specific model on pseudobulk (pb) data], respectively. Every dot is the performance (Pearson correlation) across one fold of the 20-fold CV. (E) Performance of scXpresso_{cp,pb} [cell population-specific (cp) model on pseudobulk data] summarized per tissue. Every dot represents the model's performance on a cell population in that tissue (median Pearson correlation across the 20 folds). (F) Performance of scXpresso_{cp,pb} on the different lung cell populations. The grey line indicates the median performance across all cell populations. Every dot is the performance across one fold of the 20-fold CV.

from [Mus_musculus_at_single_cell_resolution/27733](https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GSE132040). To extract input features, we downloaded the reference genome (MM10-PLUS) that was used during the alignment from: <https://s3.console.aws.amazon.com/s3/object/czb-tabula-muris-senis?region=us-west-2&prefix=reference-genome/MM10-PLUS.tgz>.

The four bulk datasets (spleen, lung, limb muscle, and bone marrow) from the Tabula Muris were downloaded from <https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GSE132040> (Schaum et al., 2019). For the bulk data, we used the same reference genome as for the single-cell data.

2.3.2 Human motor cortex data

The human motor cortex data from the Allen Institute (Bakken et al., 2021) was downloaded from the Cytosplere Comparison Viewer. We downloaded the reference genome (version GRCh38.p2) and corresponding GTF file with information about the location of transcription start sites of the genes here: https://www.gencodegenes.org/human/release_22.html.

2.4 Aggregated expression values

First, we normalized the count matrices. For the single-cell datasets, we performed library size normalization in the same

way as The Tabula Muris Consortium: i.e., counts per million for the smart-seq2 data and counts per ten thousand for the 10X data (Schaum et al., 2018). For the bulk Tabula Muris data, we performed TPM normalization. For the single-cell datasets, we used the annotations defined by the authors to aggregate the expression values per tissue or per cell population using $\log_{10}(\text{mean}(x))$ (without pseudocount) into pseudobulk values. The advantage of not adding a pseudocount is that the distribution looks more like a normal distribution, which makes it easier to train the models (Supplementary Figure S2). A limitation, however, is that we could not calculate the exact value for genes that were not expressed in any of the cells. For these genes, we replaced the pseudobulk values with -4 in the Tabula Muris and -5 in the motor cortex dataset, since this extrapolated well (Supplementary Figure S2). For the bulk data, we aggregated over the samples instead of the cells. Here, we set the genes that are not expressed in any of the samples to -4 . We standardized the expression values before running the model such that the average expression of all genes in each cell population or tissue is zero and the standard deviation is one. Before analyzing the results and comparing the predictions across cell populations, we undid the z-score normalization but kept the log normalization.

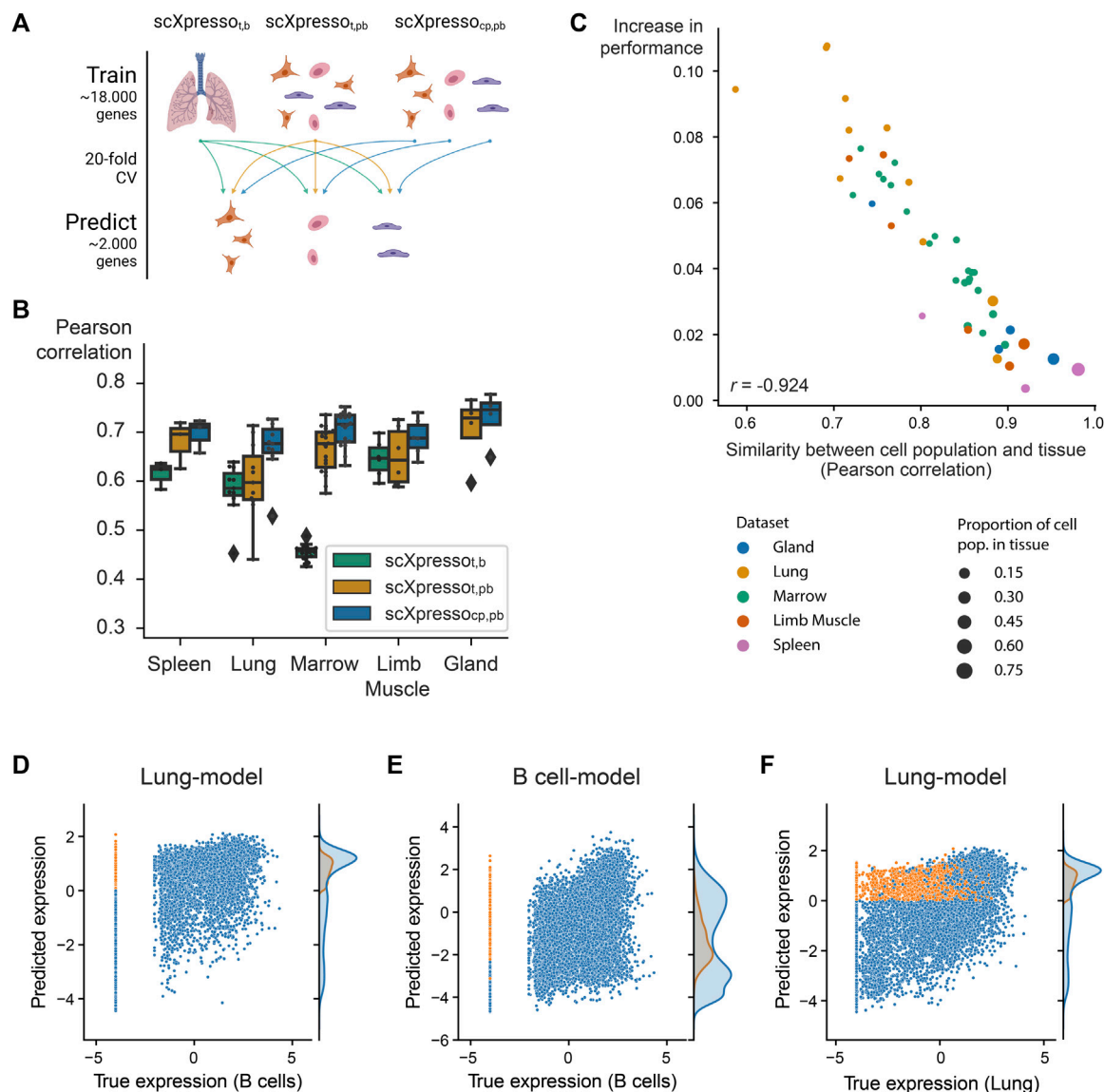


FIGURE 2

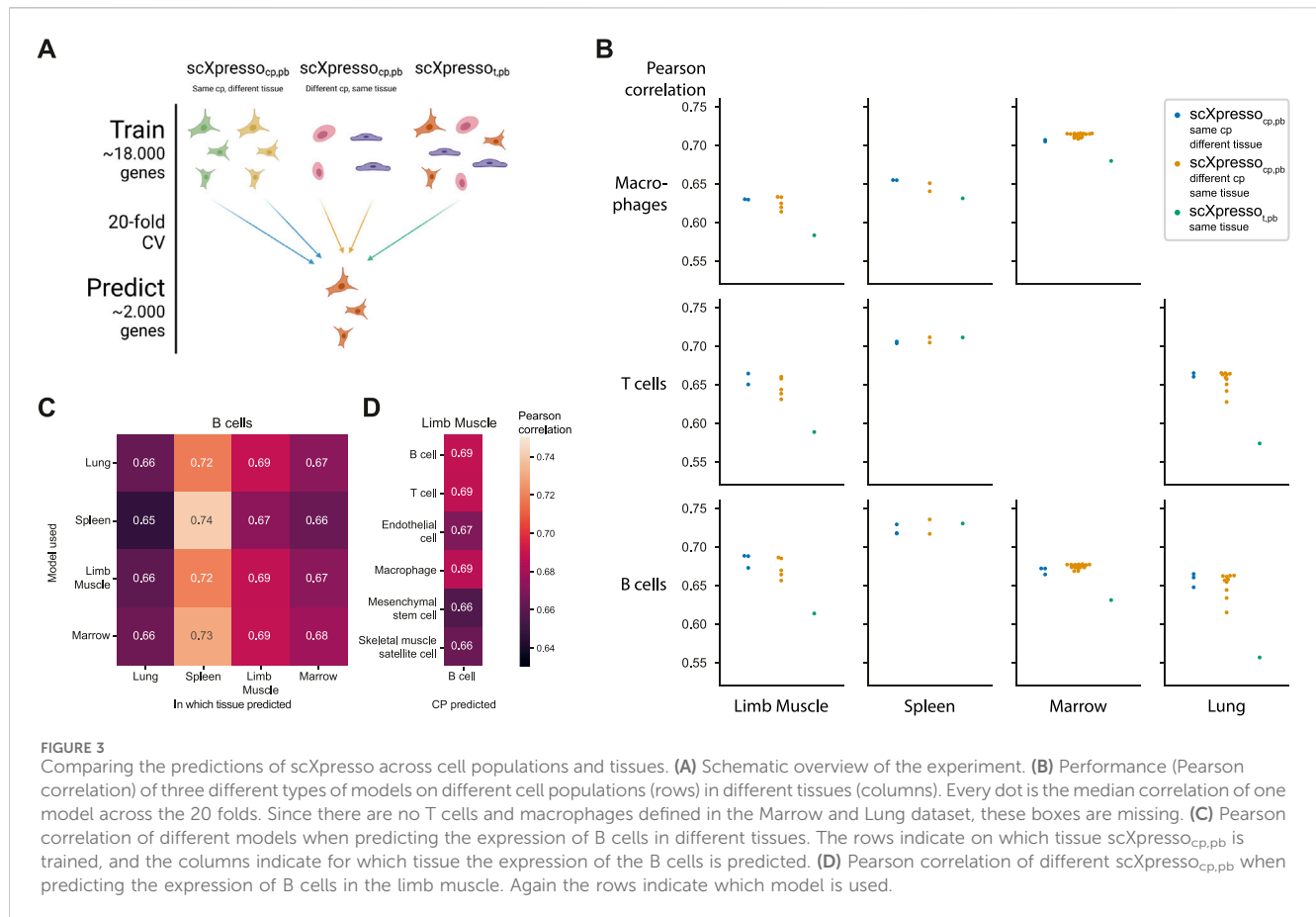
Comparison of the three scXpresso models for making cell population-specific predictions. (A) Schematic overview of the experiment. (B) Boxplot showing the performances of $scXpresso_{t,b}$ [tissue-specific (t) model on bulk (b) data], $scXpresso_{t,pb}$ [tissue-specific model on pseudobulk (pb) data], and $scXpresso_{cp,pb}$ [cell population-specific (cp) on pseudobulk (pb) data] on the cell population-specific task. Every point in the boxplot is the performance of a model on one cell population in that tissue (median Pearson correlation across the 20 folds). (C) Similarity between a cell population and corresponding tissue (Pearson correlation between the true pseudobulk expression values vs. the increase in performance ($\Delta_{cp,t}$, median Pearson correlation of $scXpresso_{cp,pb} - scXpresso_{t,pb}$). Every dot is a different cell population and the colors represent the different tissues. (D–F) Comparing the predictions made by the lung tissue model (lung-model) and the B cell population model (B cell-model). Genes where the lung-model predicts a too-high value are plotted in orange. (D,E) True expression of the B cells vs. predicted expression by the (D) lung-model and (E) B cell-model. (F) True expression of the lung cells vs. predicted expression of the lung model.

2.5 Input features

2.5.1 Sequence around the transcription start site

Before extracting the sequences around the transcription start site, we removed genes that are transgenes, ERCC spike-ins, genes without a coding region, and genes on the Y chromosome. This resulted in 20,467 mouse genes and 18,138 human genes. Some genes had multiple transcripts.

We downloaded a list with canonical transcripts for each gene from biomaRt and we used the transcript and transcription start site belonging to the canonical transcript. If the canonical transcript was not defined, we used the transcript that had the longest coding region. After having defined the transcription start site for each gene, we used seqkit (Shen et al., 2016) to extract sequences from the FASTA file containing the reference genome.



2.5.2 Half-life time features

For every gene, we extracted five half-life time features: 5' UTR length, 3' UTR length, ORF length, intron length, and exon junction density ($\frac{\#exons}{length_{ORF}} \times 1000$). We obtained these features by first filtering the GTF files for the canonical or longest transcript. The 5' UTR length is the length of the sequence from the start of the first exon to the start codon. The 3' UTR length is the length of the sequence from the last coding sequence to the end of the last exon. The ORF length is the sum of the length of the coding sequences. The intron length is the length of the transcript minus the length of the ORF, 5' UTR, and 3' UTR. All features are log-normalized using $\log_{10}(x + 0.1)$ and afterwards z-scaled.

2.6 Evaluating the predictions

For every gene in the test dataset, we averaged the predictions of the five models we trained. We evaluated the performance for every cell population by calculating the Pearson correlation between the true and predicted expression of the genes in the test set. To evaluate the increase in performance between the tissue-specific (*t*) pseudobulk (*pb*) and cell population-specific (*cp*) pseudobulk (*pb*) model on the Tabula Muris datasets, we calculate: $\Delta_{cp,t} = \text{median Pearson correlation}(\text{scEP}_{cp,pb}) - \text{median Pearson correlation}(\text{scEP}_{t,pb})$. On the motor cortex dataset, we also evaluated the

performance of each gene by calculating the Pearson correlation between the true and predicted expression per cell population.

2.7 In silico saturation mutagenesis

For *CACNA1I*, we mutated all positions *in silico*, which means we tested all possible substitutions at every position. We undid the z-score normalization and calculated the difference between the original (wild-type) prediction and the mutated prediction. The prediction models used during these experiments were the models where *CACNA1I* itself was originally in the test set. For every position, we only plotted one predicted difference in expression in Figure 4E. This is the substitution that was predicted to have the largest absolute effect. We downloaded the locations of the candidate *cis*-regulatory elements that fall within the input region for *CACNA1I* from screen registry v3 (release date 2021) (ENCODE Project Consortium et al., 2020). When plotting the difference between two cell populations, we ignored the positions where one is positive and the other predicts a negative effect. This rarely happened and if it was the case, the predicted effect was very small.

For the 2,000 highly variable genes, selected using scanpy (Wolf et al., 2018), we applied ISM similar as described for *CACNA1I*. For

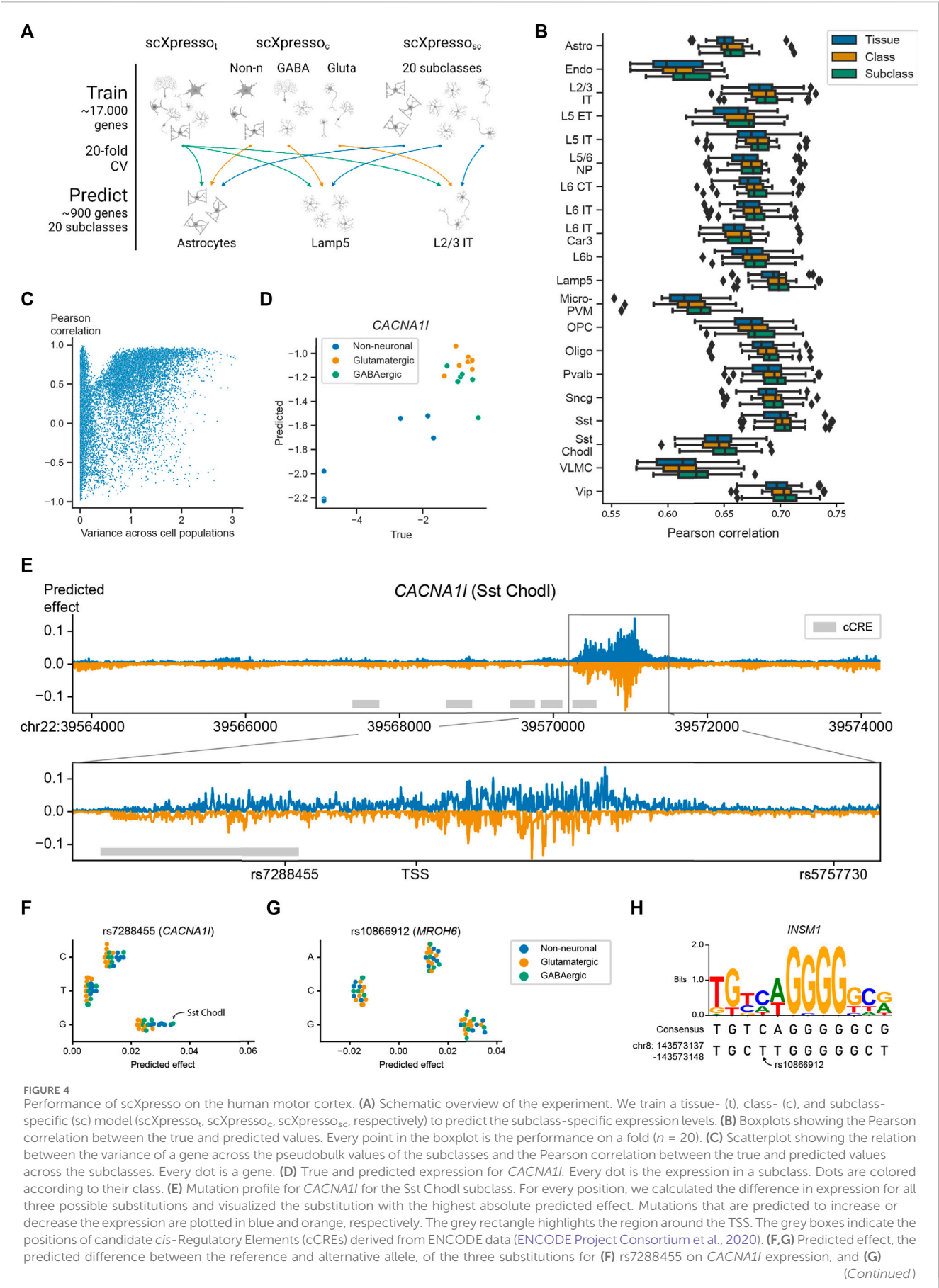


FIGURE 4 (Continued)

rs10866912 on *MROH6* expression. Every dot is one subclass and the dots are colored according to the class. (H) Sequence logo and the consensus sequence for the *INS1* transcription factor motif together with the sequence of the reference genome (bottom line).

every position we then calculated the average maximum absolute predicted effect:

$$y_{\max}(i) = \frac{1}{2000} \sum_{g \in HVG} \max_{alt \in \{A,C,G,T\}, alt \neq ref} |y_{pred,g,ref}(i) - y_{pred,g,alt}(i)|$$

where i indicates the genomic position, HVG is the list of highly variable genes, ref indicates the reference allele, and alt indicates the alternative allele.

2.8 Comparison to other models

2.8.1 Enformer

Enformer uses the DNA sequence to predict reads for 5,313 human tracks which include CAGE, DNase, CHIP, and ATAC-seq (Avsec et al., 2021a). Here, we only looked at the effect of a variant on the CAGE tracks that are related to the brain (77 tracks in total, see Supplementary Table S3). Enformer predicts the effect of variants on 128 bp bins. When predicting the effect of a variant on the CAGE reads, we looked at the effect on the bin containing the TSS.

2.8.2 ExPecto

ExPecto predicts gene expression for 218 tissues and cell lines (Zhou et al., 2018). Here, we only focused on 27 outputs that are related to the brain (Supplementary Table S4). We used the ExPecto web server to predict the effect of the variants (<https://hb.flatironinstitute.org/expecto/?tabId=3>). ExPecto is trained using Hg19 instead of Hg39. We used the R-package SNPlocs.Hsapiens.dbSNP155.GRCh37 (v 0.99.23) to lift-over the variants. Using ExPecto we could not predict the effect of all variants, since for some variants there was no location in Hg19 found, some were too far away from a TSS, and some were linked to a different gene than we were interested in (see Supplementary Table S5 for an explanation per variant).

2.8.3 Xpresso

We trained the Xpresso model on bulk RNA-seq data from the precentral gyrus (Agarwal and Shendure, 2020). The data from two individuals were downloaded from the Allen Human Brain Atlas: <https://human.brain-map.org/static/download> (H0351.2001, H0351.2002). We used the normalized matrices. Labels were created as described in the Xpresso paper: we took the median expression across the 6 precentral gyrus samples, log-normalized the output using $\log_{10}(x + 0.1)$, and z-score normalized the expression. Similar to scXpresso, we trained the model using 20-fold cross-validation. Per fold, we trained 10 runs and used the model with the lowest MSE on the validation data [as described in Agarwal and Shendure (2020)]. Afterwards, we predicted the effect of the variants. We could not predict the effect of all variants, since

some genes were not measured in the bulk RNA-seq data and for some genes, there were no Xpresso input features defined (see Supplementary Table S5 for an explanation per variant).

3 Results

3.1 Predicting cell population-specific gene expression using scXpresso

Here, we present scXpresso, a multitask convolutional neural network (CNN) to predict cell population-specific gene expression using genomic sequences only (Figure 1A; Supplementary Figure S1). We developed scXpresso by adapting the Xpresso model (Agarwal and Shendure, 2020), which was originally designed for bulk data, to single-cell data. Similar to Xpresso, we use two types of input to the model: 1) the DNA sequence around the transcription start site (TSS) (7 kb upstream—3.5 kb downstream) to model transcription, and 2) five half-life time features (5' UTR length, 3' UTR length, ORF length, intron length, and exon junction density) to model mRNA degradation. We input the one-hot encoded DNA sequence into a CNN. The output of the CNN is concatenated with the half-life time features and fed to a fully connected network (see Methods). Since our model is a multitask CNN, the desired output of the fully connected network is the gene expression for every cell population. We predict expression per cell population instead of per cell to achieve more stable predictions with less noise as single-cell data is known to be quite sparse. To obtain one expression value per cell population, we aggregate the single-cell expression into pseudobulk measurements (see Methods).

Since single-cell and bulk data have different characteristics, we tested whether scXpresso performs equally well on single-cell and bulk data. We used scRNA-seq data from five different tissues (limb muscle, spleen, gland, marrow, and lung) from the Tabula Muris (Schaum et al., 2018) (Supplementary Table S6). Here, we used cells isolated via FACS that were sequenced using the Smart-seq2 protocol. Using the annotations defined by the authors, we aggregate the values per cell population and per tissue into pseudobulk values. For four tissues (limb muscle, spleen, marrow, and lung), there are also bulk RNA-sequencing datasets available (Supplementary Table S7). We compared the pseudobulk to the bulk expression per tissue and noticed that these are indeed correlated ($r_{\text{muscle}} = 0.69$, $r_{\text{spleen}} = 0.71$, $r_{\text{marrow}} = 0.50$, $r_{\text{lung}} = 0.67$) (Supplementary Figure S3).

Next, we trained three different models: 1) a tissue-specific (t) model on the bulk (b) values (scXpresso_{t,b}), 2) a tissue-specific model on the pseudobulk (pb) values (scXpresso_{t,pb}), 3) a cell population-specific (cp) model on the pseudobulk values (scXpresso_{cp,pb}) (Figure 1B). The cell population-specific model is, in contrast to the tissue-specific models, a multitask model

that predicts the expression of all cell populations in a tissue simultaneously. We evaluated the performance of the models by calculating the Pearson correlation between the true and predicted expression values. In general, the tissue-specific models trained on pseudobulk reach higher performance than the models trained on bulk (Figures 1C, D). Even though the bulk and pseudobulk values are correlated, the pseudobulk distributions are bimodal compared to the normally distributed bulk data (Supplementary Figures S3, S4). This turns the problem more into a classification problem (is a gene low or high expressed), which might be easier to learn. On average, predicting cell population-specific expression is more difficult than predicting tissue-specific expression (Figures 1D, E): scXpresso_{cp,pb} performs slightly worse than scXpresso_{t,pb} (median correlation of 0.71 vs. 0.75), but still better than scXpresso_{t,b} (0.58).

One of the adaptations to Xpresso is that scXpresso_{cp,pb} is a multitask model. This slightly increases the performance compared to a single-task model (Supplementary Figure S5) but mainly makes the model computationally more efficient. The marrow-FACS dataset, for instance, contains 22 cell populations. Since the single-task and multitask models need the same training time (approximately 30–60 min), this gives a 22x speed up.

The Tabula Muris scRNA-seq datasets were generated using two different protocols: 10X Genomics, a droplet-based method, and FACS-based Smart-seq2, a plate-based method. When comparing scXpresso_{t,pb} and scXpresso_{cp,pb} trained on the two different protocols, e.g., lung-droplet vs. lung-FACS, we conclude that they perform equally well (Figures 1D, E; Supplementary Figures S6, S7). Depending on the tissue and cell population, one performs slightly higher than the other, but there are no significant differences. This is as expected since the pseudobulk values of both protocols are highly correlated (Pearson correlation > 0.85) (Supplementary Figure S8). Hence, the protocol used to create the single-cell dataset does not influence the results.

For scXpresso_{cp,pb}, we tested how the two types of input features, DNA sequence and half-life time, influence the performance. We tested different lengths of the input sequence and whether one of the two features was enough to predict expression (Supplementary Figure S9). A range of different sequence lengths results in the same performance (3.5–3.5, 7–3.5, and 10–5 kb upstream-downstream). A longer sequence gives more information but also adds more noise. Since the model also becomes more complex, more parameters have to be learned and it takes more time and memory to train the model. Therefore, we decided to use 7 kb upstream and 3.5 kb downstream for further experiments. We also observed that adding the half-life time features results in higher performance, suggesting that these features are not easily captured from DNA sequences alone.

For the cell population-specific models, the performance varies considerably across different populations (Figure 1E). Comparing the populations in the lung dataset, for instance, the performance of the endothelial cells is very high compared to leukocytes (Figure 1F; Supplementary Figure S10). In general, the performance of scXpresso increases when more genes and cells are measured in a population (Figure 1F; Supplementary Figure S11). The leukocyte population is small (35 cells) and fewer genes are non-zero compared to other cell populations in the lung (8,678 out of 20,467 vs. 12,715 on average). The ciliated cell population, on the other hand, is also small (25 cells), but this model reaches a higher

performance. In this cell population, however, more genes were non-zero (11,717) compared to the leukocyte population. Hence, to train the model, we need a good representation of the cell population that includes enough expressed genes.

In all previous experiments, we evaluated scXpresso using 20-fold cross-validation with the genes randomly divided over the folds. The results could be positively biased if genes from the same chromosome are in different folds. Therefore, we also evaluated the models using cross-chromosomal cross-validation. This slightly reduces the models' performance, but the difference is not significant (lowest *p*-value = 0.11 for myeloid cells, two-sample Wilcoxon rank sum test) (Supplementary Figure S12).

3.2 Cell population-specific models outperform tissue-specific models

Now that we know that all models are well-trained, we predicted cell population-specific expression using the three different models to see whether increasing the resolution of the models increases the performance (Figure 2A). Since scXpresso_{t,b} and scXpresso_{t,pb} were trained using tissue-specific expression values, these models predict the same value for every cell population. On all datasets, scXpresso_{cp,pb} outperforms the tissue-specific models, which shows the benefit of training the models on a higher resolution (Figure 2B; Supplementary Figure S13A). Especially in more heterogeneous tissues, where the gene expression of cell populations is weakly correlated to the corresponding tissue, we see a large improvement (Figure 2C; Supplementary Figure S13B). For the lung-FACS dataset, for instance, the performance increases the most for immune cell populations ($\Delta_{cp,t}$ for B cells: 0.11, NK cells: 0.11, T cells: 0.09; see **Methods**) and the least for lung-specific populations ($\Delta_{cp,t}$ for stromal cells: 0.01, endothelial cells: 0.03, epithelial cells: 0.05). In the B cells in the lung, 4,081 genes are not expressed in any of the cells and thus have a log-normalized expression of -4 , but for which the tissue-specific model predicts a positive log-normalized expression value (Figure 2D). In contrast, the model trained on B cells predicts a lower expression for these genes (Figure 2E). Almost all these genes, however, are expressed in the lung (in the non-B cells), the lung-model learned this correctly too (Figure 2F).

Some of the Tabula Muris datasets contain similar cell populations. For instance, B cells, macrophages, and T cells are measured in four, three, and three tissues, respectively. We hypothesized that if our models are cell population-specific, they should accurately predict the expression of a cell population in one tissue with a model trained on the same cell population but from another tissue [even though a cell's tissue will slightly change the expression for (some) genes]. To test this, we predicted the expression for common cell populations using three different types of models: 1) scXpresso_{cp,pb} trained on the same cell population, but from a different tissue, 2) scXpresso_{cp,pb} trained on a different cell population, but from the same tissue, 3) scXpresso_{t,pb} trained on the same tissue (Figure 3A). For example, we predict the expression of B-cells in the limb muscle, using 1) a model trained on B-cells in the lung, 2) a model trained on endothelial cells in the limb muscle, and 3) a model trained on the limb muscle. Again, the cell population-specific models outperform

the tissue-specific models, even though they predict either a different dataset or a different cell population than they were trained on (Figure 3B; Supplementary Figures S14, S15). This indicates that if you want to train a model for a cell population from a specific tissue where no single-cell data is available, you are better off using a model trained on a similar cell population from a different tissue than relying on a tissue-specific model. Whether a model trained on a different cell population and the same tissue performs better than a model trained on the same cell population but a different tissue, differs per tissue and cell population. For example, when predicting the expression of B cells in the limb muscle, the models trained on B cells in the marrow and lung even outperform the model trained on B cells in the limb muscle itself (Figure 3C). But, the models trained on different cell populations within the limb muscle perform variably when predicting B cells (Figure 3D). The models trained on immune populations, e.g., T cells or macrophages, perform similarly, but the muscle-specific populations perform worse. This difference between the B cell and the endothelial, mesenchymal stem cell, and skeletal muscle satellite cell models might seem small but is significant across the 20 folds [p -value = $9.5e-07$ for all three populations, one-sided Wilcoxon rank sum test (Mann and Whitney, 1947; Demšar, 2006)]. Even though the differences are small, this indicates that our models indeed learn cell population-specific features.

3.3 scXpresso learns expression patterns across human brain cell populations

Next, we applied scXpresso to a human brain dataset of the motor cortex (Bakken et al., 2021). This dataset is annotated at different resolutions including a class (GABAergic, glutamatergic, and non-neuronal) and subclass (20 subclasses) level. Again, we trained models of different resolutions: a tissue- (t), class- (c), and subclass-specific (sc) model (scXpresso_t, scXpresso_c, and scXpresso_{sc}, respectively). We used the trained models to predict the subclass-specific expression values (Figure 4A). Since scXpresso_t was trained on the tissue-specific pseudobulk expression, it predicts the same expression for all subclasses. The class-specific model, on the contrary, is a multitask model. Here, we use the predictions of the parent class to predict the expression of each subclass (i.e., subclasses belonging to the same parent class are predicted to have the same expression) (Supplementary Figure S16). Similar to the Tabula Muris, we observed that increasing the resolution increases the performance: scXpresso_{sc} outperforms scXpresso_c, which outperforms scXpresso_t (Figure 4B). For some subclasses, e.g., L2/3 IT, the performance barely improves when comparing scXpresso_{sc} with scXpresso_c, which happens when the true expression values of the subclass and corresponding class are strongly correlated, similar as for the Tabula Muris case (Supplementary Figure S17).

Since genes with variable expression across subclasses are often interesting to study, we tested whether scXpresso_{sc} can learn the correct pattern for a gene across the subclasses. For every gene, we calculate the Pearson correlation between the true and predicted expression across the subclasses. If the expression of a gene shows some variance across the subclasses, scXpresso_{sc} predicts

the pattern correctly (Figure 4C). An example is *CACNA1I*, a gene coding for a subtype of voltage-gated calcium channel that has been associated with schizophrenia (Li et al., 2017; Pardiñas et al., 2018; Ikeda et al., 2019; Lam et al., 2019; Yao et al., 2021). Here scXpresso_{sc} correctly learns that the expression in neuronal populations is higher than in non-neuronal ($r = 0.90$) (Figure 4D).

3.4 *In silico* saturation mutagenesis reveals the most interesting GWAS variants

Since scXpresso can predict expression from the DNA sequence, we expect that it can also predict how the expression changes when the sequence is mutated. Therefore, we applied *in silico* saturation mutagenesis (ISM) to the sequence of *CACNA1I* and evaluated the predicted change in gene expression (Zhou and Troyanskaya, 2015; Kelley et al., 2016; Kelley et al., 2018; Avsec et al., 2021a). When comparing scXpresso_{sc} predictions for the Sst Chodl subclass across all possible mutations, we find mutations in the region around the TSS to affect the expression of the *CACNA1I* gene the most (Figure 4E). When applying ISM to the 2,000 highly variable genes in the data, the maximum absolute predicted effect is highest around the TSS as well (Supplementary Figure S18). Note, that we did not use the TSS location as input to the model, consequently, the model correctly identified that this is the most important region for transcriptional regulation. No other regions within our input window were found that affect the expression that strongly.

Besides visualizing the mutation pattern for one subclass, we can also visualize how ISM affects two subclasses differently. As an example, we compared the scXpresso_{sc} predictions for the Sst Chodl subclass and the L2/3 IT subclass (Supplementary Figure S19). These predictions show that the Sst Chodl subclass is more sensitive to mutations than the L2/3 IT class for *CACNA1I*, which might be explained by the fact that *CACNA1I* is also higher expressed in Sst Chodl cells.

In addition, ISM can be used to prioritize variants of interest for diseases. *CACNA1I* is linked to 18 Schizophrenia-associated variants according to the NHGRI-EBI Catalog (Buniello et al., 2019). Two of these variants, rs7288455 and rs5757730, fall within our input region (7 kb upstream and 3.5 kb downstream of the *CACNA1I* TSS). Mutating the reference A allele with the C or G variant at the position of rs7288455 increases the predicted expression for all cell populations (Figure 4F). The disease-associated variant, the A allele, is expected to decrease the expression (Buniello et al., 2019; Yao et al., 2021), which is in line with our predictions, although it is not known whether this is subclass-related. Our model suggests that the expression of *CACNA1I* increases the most in the Sst Chodl subclass. Interestingly, for the Sst Chodl subclass, this mutation results in one of the largest differences in *CACNA1I* expression amongst all other induced mutations (top 1% mutations with the strongest effect) (Supplementary Figure S20). For the other variant, rs5757730, which lies in an intronic region, we see no difference in expression (Supplementary Figure S21). Further supporting our predictions, rs7288455, but not rs5757730, overlaps with an ENCODE candidate *cis*-regulatory element. These results show that scXpresso can be used to prioritize GWAS hits.

In total, there are 3,971 GWAS variants associated with Schizophrenia in the NHGRI-EBI Catalog (Buniello et al., 2019). We focused on those genes that have two or more variants in the input region (20 genes, 49 variants) (Supplementary Table S5). For these variants, we predicted the effect of all possible substitutions to prioritize the likely causal variants (Supplementary Figure S22). For most genes, scXpresso predicts a profound effect for only one of the variants. For instance, when substituting “A” with “C” for the *HLA-B* variant rs2507989, the predicted expression of *HLA-B* decreases, while none of the mutations at the other variant positions of *HLA-B*, i.e., rs139099016 and rs1131275, are predicted to affect the expression. Noteworthy, rs1131275 is classified as a missense variant and thus not expected to alter transcription (Buniello et al., 2019). For some genes, however, all variants seem to barely affect the expression.

Next, we checked if we could interpret the model predictions by characterizing the genomic sequences identified by scXpresso to have a strong effect on gene expression. For the *MROH-6* variant rs10866912, two substitutions are predicted to create an opposite effect. Substituting the reference “T” with a “C” is predicted to decrease the expression while mutating with a “G” is predicted to increase the expression (Figure 4G). This variant is part of a binding site for the transcription factor *INSM1*, a transcriptional repressor (Castro-Mondragon et al., 2022) (Figure 4H). When substituting the “T” with a “C”, the sequence of the reference genome becomes more similar to the consensus motif, while substituting with a “G” makes the two sequences more dissimilar. This supports the predictions from scXpresso.

We compared our scXpresso predictions for these Schizophrenia variants to the predictions of Enformer, ExPecto, and Xpresso. For Enformer and ExPecto we used their pre-trained models which predict the expression for 5,313 and 218 tissues/cell lines, respectively. Here, we only focused on the predictions related to the healthy brain (77 tracks for Enformer, 27 for ExPecto). For Xpresso, there were no pre-trained models for the brain available, so we trained the Xpresso model ourselves using bulk RNA-seq samples from the precentral gyrus, which is the region containing the motor cortex (see Methods). The expression values of the precentral gyrus are correlated to the pseudobulk expression values of the motor cortex (Supplementary Figure S23A, $r = 0.68$). Similar to scXpresso, we used a 20-fold cross-validation to train the Xpresso model. The model is well-trained and reached a similar median correlation on the precentral gyrus as the scXpresso models on the motor cortex subclasses (Figure 4B, S23B-C, $r = 0.69$). Supplementary Figure S24 shows the predictions for all models for the variants related to Schizophrenia. Using Xpresso and ExPecto we could not predict the effect of all variants, since some genes were missing from the data and some variants were lost during conversion from Hg38 to Hg19 (Supplementary Table S5) (see Methods). It is challenging to compare the predictions of the different methods since all models are trained on different brain regions or cell lines. Enformer usually predicts the same effect for the three different possible nucleotide mutations, e.g., for rs1131275 it predicts that all three substitutions decrease the expression. This variant, however, is classified as a missense variant, so we do not expect it to alter transcription (Buniello et al., 2019). For rs7288455, the variant close to *CACNA1I*, both scXpresso and Xpresso predict a similar effect, while Enformer and ExPecto predict only a very

minimal effect. For rs10866912, the variant close to *MROH-6*, we showed that scXpresso could learn the TF binding site of *INSM1* while all the other models miss this pattern. These results overall illustrate the benefit of training prediction models on single-cell data.

4 Discussion

We presented scXpresso, a model to predict cell population-specific gene expression using the genomic sequence. We showed that scXpresso outperforms tissue-specific bulk and pseudobulk models especially when the expression profile of a cell population is dissimilar to that of the corresponding tissue. All scXpresso models reach a Pearson correlation of approximately 0.7 regardless of the cell population or tissue trained on. Additionally, the model learned the importance of the region around the TSS, transcription factor binding motifs (such as for *INSM1*), and the expression pattern of genes across different cell populations. Together, our findings show the potential of using single-cell data for predicting gene expression from sequence information in complex heterogeneous tissues.

We showed that it is possible to prioritize GWAS variants using scXpresso. Considering the expression of *CACNA1I*, we noticed that one variant, which overlaps with an ENCODE *cis*-regulatory element, is predicted to have a large effect, while another variant was predicted to have a negligible effect. The latter could be because the variant might affect splicing (which our model does not differentiate), the variant could be in a linkage disequilibrium block with other (associating) variants, or the variant could affect a more distant gene.

Comparing the predicted effects for mutations by scXpresso to other sequence-to-expression prediction models quantitatively is difficult as the true effect of these variants on specific brain regions and/or cell populations is unknown. We have shown that for a previously identified variant close to *CACNA1I* gene, both Xpresso and scXpresso predict an increase in expression, while ExPecto and Enformer predict a marginal effect. Note that, ExPecto and Enformer are not trained on specific brain regions, or cell population-specific data, but contain bigger structures such as the frontal cortex or frontal lobe. Hence, these models miss the cell population-specific effect of this variant. Training these models on cell population-specific scRNA data could be an interesting next step.

Using our model, it is not possible to test trans-effects of variants as our model uses a limited genomic sequence region as input. Consequently, we could only test two variants related to Schizophrenia for *CACNA1I*, out of the 18 variants associated with *CACNA1I* (Buniello et al., 2019). Ideally, we would increase the length of the input sequence, however, it is not easy to learn long-range interactions using CNNs. The Enformer model, which uses a 200 kb sequence as input, tackles this problem by combining transformers and CNNs (Avsec et al., 2021a). Unfortunately, the Enformer model predicts CAGE reads instead of expression values, so we cannot trivially extend it or use it for single-cell data. An alternative approach might be to use their well-trained model to get an embedding for every input sequence and use this embedding to predict cell population-specific expressions.

We input the DNA sequence and five half-life time features to scXpresso. However, certain transcript features, which are related to the half-life time features, can predict zeros in the scRNA-seq data (Lipnitskaya et al., 2022). Whether the observed zeros in scRNA-seq data are technical artifacts or biologically informative is an ongoing debate. We believe that the zeros are biologically informative since binarized data can be used for downstream analysis, resulting in comparable results to those obtained using scRNA-seq counts (Bouland et al., 2023). Furthermore, we would like to highlight that the performance of the cell population-specific pseudobulk models when trained on sequence-only information is also not much lower as compared to both sequence and half-life time features (Supplementary Figure S9). This observation supports our conclusion that the half-life time features are not biasing the models towards scRNA-seq artifacts.

Two future enhancements that we envision that could improve the performance of our model are related to the half-life time features and the output of the model. Currently, we extract five features from the mRNA sequence to approximate the half-life time. Recently, a new model, Saluki, was developed that could predict mRNA degradation rates directly from the sequence of the gene (Agarwal and Kelley, 2022). Replacing the currently used features with those predicted by the Saluki model, or combining these features, might improve the cell population-specific predictions. A second potential improvement relates to the current output of scXpresso, which is the pseudobulk expression for every cell population, i.e., the average gene expression across all cells from that population. However, this ignores the variance within the population. It might be more beneficial to predict the distribution of gene expression across each population, instead of just one aggregated value.

In summary, we have shown the potential of predicting cell population-specific gene expression from genomic sequences by leveraging the resolution of single-cell data, opening the way for many new developments in this area.

Data availability statement

Publicly available datasets were analyzed in this study. This data can be found here: Zenodo (pseudobulk expression values, trained models, predictions): <https://doi.org/10.5281/zenodo.7044908>, TM data single-cell: https://figshare.com/projects/Tabula_Muris_Transcriptomic_characterization_of_20_organisms_and_tissues_from_Mus_musculus_at_single_cell_resolution/27733, TM bulk: <https://www.ncbi.nlm.nih.gov/geo/query/acc.cgi?acc=GSE132040>, TM genome: <https://s3.console.aws.amazon.com/s3/object/czb-tabula-muris-senis?region=us-west-2&prefix=reference-genome/MM10-PLUS.tgz>, M1 single-cell: cytoplasm comparison viewer M1 gtf file: https://www.encodegenes.org/human/release_22.html. The code to reproduce the figures, train your own models, show the effect of variants, and do in silico saturation mutagenesis can be found on GitHub: <https://github.com/lcmmichielsen/scXpresso>.

References

Agarwal, V., and Kelley, D. R. (2022). The genetic and biochemical determinants of mRNA degradation rates in mammals. *Genome Biol.* 23, 245. doi:10.1186/s13059-022-02811-x

Author contributions

LM: Conceptualization, Formal Analysis, Methodology, Software, Visualization, Writing—original draft, Writing—review and editing. MR: Conceptualization, Funding acquisition, Supervision, Writing—review and editing. AM: Conceptualization, Funding acquisition, Supervision, Writing—review and editing.

Funding

The authors declare financial support was received for the research, authorship, and/or publication of this article. This research was supported by an NWO Gravitation project: BRAINSCAPES: A Roadmap from Neurogenetics to Neurobiology (NWO: 024.004.012).

Acknowledgments

We would like to thank Dr. Stavros Makrodimitris for his insightful discussion and his example of PyTorch code for convolutional neural networks. Figures 1B, 2A, 3A, 4A and Supplementary Figure S16 were created with BioRender.com.

Conflict of interest

The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

The authors declared that they were an editorial board member of Frontiers, at the time of submission. This had no impact on the peer review process and the final decision.

Publisher's note

All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers. Any product that may be evaluated in this article, or claim that may be made by its manufacturer, is not guaranteed or endorsed by the publisher.

Supplementary material

The Supplementary Material for this article can be found online at: <https://www.frontiersin.org/articles/10.3389/fbinf.2024.1347276/full#supplementary-material>

Agarwal, V., and Shendure, J. (2020). Predicting mRNA abundance directly from genomic sequence using deep convolutional neural networks. *Cell Rep.* 31, 107663. doi:10.1016/j.celrep.2020.107663

- Avsec, Ž., Agarwal, V., Visentin, D., Ledsam, J. R., Grabska-Barwinska, A., Taylor, K. R., et al. (2021a). Effective gene expression prediction from sequence by integrating long-range interactions. *Nat. Methods* 18, 1196–1203. doi:10.1038/s41592-021-01252-x
- Avsec, Ž., Weilert, M., Shrikumar, A., Krueger, S., Alexandari, A., Dalal, K., et al. (2021b). Base-resolution models of transcription-factor binding reveal soft motif syntax. *Nat. Genet.* 53, 354–366. doi:10.1038/s41588-021-00782-6
- Bakken, T. E., Jorstad, N. L., Hu, Q., Lake, B. B., Tian, W., Kalmbach, B. E., et al. (2021). Comparative cellular analysis of motor cortex in human, marmoset and mouse. *Nature* 598, 111–119. doi:10.1038/s41586-021-03465-8
- Boulant, G. A., Mahfouz, A., and Reinders, M. J. T. (2023). Consequences and opportunities arising due to sparser single-cell RNA-seq datasets. *Genome Biol.* 24, 86. doi:10.1186/s13059-023-02933-w
- Buniello, A., MacArthur, J. A. L., Cerezo, M., Harris, L. W., Hayhurst, J., Malangone, C., 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. doi:10.1093/nar/gky1120
- Castro-Mondragon, J. A., Riudavets-Puig, R., Rauluseviciute, I., Lemma, R. B., Turchi, L., Blanc-Mathieu, R., et al. (2022). JASPAR 2022: the 9th release of the open-access database of transcription factor binding profiles. *Nucleic Acids Res.* 50, D165–D173. doi:10.1093/nar/gkab1113
- Demšar, J. (2006). Statistical comparisons of classifiers over multiple data sets. *J. Mach. Learn. Res.* 7, 1–30.
- ENCODE Project Consortium, Moore, J. E., Purcaro, M. J., Pratt, H. E., Epstein, C. B., Shores, N., et al. (2020). Expanded encyclopaedias of DNA elements in the human and mouse genomes. *Nature* 583, 699–710. doi:10.1038/s41586-020-2493-4
- Ikeda, M., Takahashi, A., Kamatani, Y., Momozawa, Y., Saito, T., Kondo, K., et al. (2019). Genome-wide association study detected novel susceptibility genes for schizophrenia and shared trans-populations/diseases genetic effect. *Schizophr. Bull.* 45, 824–834. doi:10.1093/schbul/sby140
- Janssens, J., Aibar, S., Taskiran, I. I., Ismail, J. N., Gomez, A. E., Aughey, G., et al. (2022). Decoding gene regulation in the fly brain. *Nature* 601, 630–636. doi:10.1038/s41586-021-04262-z
- Kelley, D. R., Reshef, Y. A., Bileschi, M., Belanger, D., McLean, C. Y., and Snoek, J. (2018). Sequential regulatory activity prediction across chromosomes with convolutional neural networks. *Genome Res.* 28, 739–750. doi:10.1101/gr.227819.117
- Kelley, D. R., Snoek, J., and Rinn, J. L. (2016). Basset: learning the regulatory code of the accessible genome with deep convolutional neural networks. *Genome Res.* 26, 990–999. doi:10.1101/gr.200535.115
- Lam, M., Chen, C.-Y., Li, Z., Martin, A. R., Bryois, J., Ma, X., et al. (2019). Comparative genetic architectures of schizophrenia in East Asian and European populations. *Nat. Genet.* 51, 1670–1678. doi:10.1038/s41588-019-0512-x
- Lambert, S. A., Jolma, A., Campitelli, L. F., Das, P. K., Yin, Y., Albu, M., et al. (2018). The human transcription factors. *Cell* 172, 650–665. doi:10.1016/j.cell.2018.01.029
- Li, Z., Chen, J., Yu, H., He, L., Xu, Y., Zhang, D., et al. (2017). Genome-wide association analysis identifies 30 new susceptibility loci for schizophrenia. *Nat. Genet.* 49, 1576–1583. doi:10.1038/ng.3973
- Lipnitskaya, S., Shen, Y., Legewie, S., Klein, H., and Becker, K. (2022). Machine learning-assisted identification of factors contributing to the technical variability between bulk and single-cell RNA-seq experiments. *bioRxiv*. doi:10.1101/2022.01.06.474932
- Mann, H. B., and Whitney, D. R. (1947). On a test of whether one of two random variables is stochastically larger than the other. *aoms* 18, 50–60. doi:10.1214/aoms/1177730491
- Nott, A., Holtman, I. R., Coufal, N. G., Schlachetzki, J. C. M., Yu, M., Hu, R., et al. (2019). Brain cell type-specific enhancer-promoter interactome maps and disease-risk association. *Science* 366, 1134–1139. doi:10.1126/science.aay0793
- Pardiñas, A. F., Holmans, P., Pocklington, A. J., Escott-Price, V., Ripke, S., Carrera, N., et al. (2018). Common schizophrenia alleles are enriched in mutation-intolerant genes and in regions under strong background selection. *Nat. Genet.* 50, 381–389. doi:10.1038/s41588-018-0059-2
- Paszke, A., Gross, S., Massa, F., Lerer, A., Bradbury, J., Chanan, G., et al. (2019). “PyTorch: an imperative style, high-performance deep learning library,” in *Advances in neural information processing systems* 32. Editors H. Wallach, H. Larochelle, A. Beygelzimer, F. d\textquotesingle Alché-Buc, E. Fox, and R. Garnett (New York: Curran Associates, Inc.), 8024–8035.
- Schaum, N., Karkanas, J., Neff, N. F., May, A. P., Quake, S. R., Wyss-Coray, T., et al. (2018). Single-cell transcriptomics of 20 mouse organs creates a Tabula Muris. *Nature* 562, 367–372. doi:10.1038/s41586-018-0590-4
- Schaum, N., Lehallier, B., Hahn, O., Hosseinzadeh, S., Lee, S. E., Sit, R., et al. (2019). The murine transcriptome reveals global aging nodes with organ-specific phase and amplitude. *bioRxiv*, 662254. doi:10.1101/662254
- Schizophrenia Working Group of the Psychiatric Genomics Consortium (2014). Biological insights from 108 schizophrenia-associated genetic loci. *Nature* 511, 421–427. doi:10.1038/nature13595
- Sharova, L. V., Sharov, A. A., Nedorezov, T., Piao, Y., Shaik, N., and Ko, M. S. H. (2009). Database for mRNA half-life of 19 977 genes obtained by DNA microarray analysis of pluripotent and differentiating mouse embryonic stem cells. *DNA Res.* 16, 45–58. doi:10.1093/dnares/dsn030
- Shen, W., Le, S., Li, Y., and Hu, F. (2016). SeqKit: a cross-platform and ultrafast toolkit for FASTA/Q file manipulation. *PLoS One* 11, e0163962. doi:10.1371/journal.pone.0163962
- Spies, N., Burge, C. B., and Bartel, D. P. (2013). 3' UTR-isoform choice has limited influence on the stability and translational efficiency of most mRNAs in mouse fibroblasts. *Genome Res.* 23, 2078–2090. doi:10.1101/gr.156919.113
- Tasic, B., Menon, V., Nguyen, T. N., Kim, T. K., Jarsky, T., Yao, Z., et al. (2016). Adult mouse cortical cell taxonomy revealed by single cell transcriptomics. *Nat. Neurosci.* 19, 335–346. doi:10.1038/nn.4216
- Tasic, B., Yao, Z., Graybiel, L. T., Smith, K. A., Nguyen, T. N., Bertagnolli, D., et al. (2018). Shared and distinct transcriptomic cell types across neocortical areas. *Nature* 563, 72–78. doi:10.1038/s41586-018-0654-5
- Vaquerez, J. M., Kummerfeld, S. K., Teichmann, S. A., and Luscombe, N. M. (2009). A census of human transcription factors: function, expression and evolution. *Nat. Rev. Genet.* 10, 252–263. doi:10.1038/nrg2538
- Wesolowska-Andersen, A., Zhuo Yu, G., Nylander, V., Abaitua, F., Thurner, M., Torres, J. M., et al. (2020). Deep learning models predict regulatory variants in pancreatic islets and refine type 2 diabetes association signals. *Elife* 9, e51503. doi:10.7554/eLife.51503
- Wightman, D. P., Jansen, I. E., Savage, J. E., Shadrin, A. A., Bahrami, S., Holland, D., et al. (2021). A genome-wide association study with 1,126,563 individuals identifies new risk loci for Alzheimer's disease. *Nat. Genet.* 53, 1276–1282. doi:10.1038/s41588-021-00921-z
- Wolf, F. A., Angerer, P., and Theis, F. J. (2018). SCANPY: large-scale single-cell gene expression data analysis. *Genome Biol.* 19, 15. doi:10.1186/s13059-017-1382-0
- Yao, X., Glessner, J. T., Li, J., Qi, X., Hou, X., Zhu, C., et al. (2021). Integrative analysis of genome-wide association studies identifies novel loci associated with neuropsychiatric disorders. *Transl. Psychiatry* 11, 69. doi:10.1038/s41398-020-01195-5
- Zhang, Y., Zhou, X., and Cai, X. (2020). Predicting gene expression from DNA sequence using residual neural network. *bioRxiv*.
- Zhou, J., Theesfeld, C. L., Yao, K., Chen, K. M., Wong, A. K., and Troyanskaya, O. G. (2018). Deep learning sequence-based *ab initio* prediction of variant effects on expression and disease risk. *Nat. Genet.* 50, 1171–1179. doi:10.1038/s41588-018-0160-6
- Zhou, J., and Troyanskaya, O. G. (2015). Predicting effects of noncoding variants with deep learning-based sequence model. *Nat. Methods* 12, 931–934. doi:10.1038/nmeth.3547