

Interpreting single-cell and spatial omics data using deep neural network training dynamics

In the format provided by the
authors and unedited

A Supplementary Material

A.1 Small intestine data analysis

We analyzed a single-cell RNA-sequencing dataset collected from mice small intestine[21]. In the original study, cells were labeled according to several annotation classes: cell type (enterocytes, goblet cells, enteroendocrine cells, tuft cells, paneth cells, transit amplifying cells, and stem cells), spatial segment (30 subsegments along the proximal to distal axis), and cell cycle phase (G1/G2M/S). We further used the spatial coarse division of the segments into five domains, as in [21] (A-[1-2], B-[3-9], C-[10-17], D-[18-24], E-[25-30]). The cell-cell variance of both confidence and variability assigned using Annotatability (Methods 2.8) were highest for the spatial annotation class (variance of confidence/variability: 0.0381/0.0029; Extended Data Figure1b), followed by cell cycle annotations (variance of confidence/variability: 0.0316/0.0010; Extended Data Figure 1a), and cell type annotations (variance of confidence/variability: 0.0290/0.0012; Extended Data Figure 1c). This result is consistent with the underlying continuous structure of the spatial signal, as many of the highly regionalized genes[21] exhibit gradient-like patterns along the proximal-to-distal axis, rather than acting as localized markers of specific spatial domains. This contrasts, for example, cell type annotations, which have known markers for each cell type.

Across the different annotation classes, some labels are harder-to-learn than others, hinting at inherent differences in the underlying biological characteristics of cells corresponding to these labels. For the spatial segmentation annotation class, spatial boundary clusters (A and E) are easier to classify (mean confidence values: 0.746 and 0.760, respectively), due to their distinct features, compared to inner clusters (B, C, and D; mean confidence values: 0.613, 0.555 and 0.599, respectively), due to their overlapping characteristics.

For the cell type annotation class, early stages of development tend to have lower confidence scores; Specifically, stem cells and transit-amplifying cells exhibit lower confidence (mean confidence: 0.785 and 0.785, respectively) compared to other cell types (mean confidence: 0.940) (Extended Data Figure 1d-f). Finally, we will focus on the cell cycle annotation class. The cell cycle is a dynamic process consisting of several phases. Even though our input labels are discrete (G1/G2M/S), the underlying trajectory cells go through has continuous features[17]. We observed that cells that have exited the dynamic process of the cell cycle and have reached their differentiated state are easier-to-learn compared to cycling cells (mean confidence scores – enterocytes: 0.955, goblet cells: 0.936, enteroendocrine cells: 0.934, tuft cells: 0.956, paneth cells: 0.848, transit amplifying cells: 0.858, and stem cells: 0.798).

A.2 PBMC dataset

A.2.1 Analysis of NK cells within the PBMC dataset

Peripheral blood NK cells are divided into two main subtypes: $CD56^{dim}$ and $CD56^{bright}$, comprising the majority and minority of peripheral NK cells, respectively[2, 3]. Within the PBMC dataset [5, 6], we found that NK cells classified by Annotatability as correctly annotated express, on average, eight known $CD56^{dim}$ marker genes[3]. Conversely, NK cells classified as ambiguously annotated express, on average, all four known $CD56^{bright}$ marker genes[3] that were not filtered out during preprocessing (Methods 2.7; Figure 2g, Supplementary Figure 3). Indeed, $CD56^{bright}$ cells belong to an earlier and less differentiated developmental stage than $CD56^{dim}$ cells[3]. Additionally, the gene expression profiles of $CD56^{bright}$ partially resemble those of T cells (Supplementary figure 2c), and specifically CD4+ naive T cells, which share some $CD56^{bright}$ marker genes (*IL7R*, *CCR7*[8]; Supplementary Figure 2a); this similarity may hinder the neural network’s ability to classify $CD56^{bright}$ cells as NK cells.

A.2.2 Analysis of monocytes within the PBMC dataset

CD14+ monocytes classified by Annotatability as ambiguously annotated express lower levels of classical monocyte markers and higher levels of intermediate monocyte markers than CD14+ monocytes classified as correctly annotated (classical monocyte markers: mean expression of *S100A12/ALOX5AP* is 2.463/0.298 in ambiguous compared to 3.850/0.685 in correctly annotated, p-value of 3.087e-83/1.23e-20 in Wilcoxon test; intermediate monocyte markers: mean expression of *GPR35/MARCO* is 0.149/0.413 in ambiguous compared to 0.105/0.202 in correctly annotated, p-value of 0.069/5.34e-08 in Wilcoxon test) (Figure 2d, h, Supplementary Figure 3). Similarly, FCGR3A+ monocytes classified as ambiguously annotated express lower levels of non-classical monocyte markers and higher levels of intermediate

monocyte markers than FCGR3A+ monocytes classified as correctly annotated (non-classical monocyte markers: mean expression of *ICAM4*/*CD79b* is 0.118/0.926 in ambiguous compared to 0.655/1.620 in correctly annotated, p-value of 2.23e-11/1.40e-06 in Wilcoxon test; intermediate monocyte markers: mean expression of *GPR35*/*MARCO* is 0.168/0.502 in ambiguous compared to 0.083/0.167 in correctly annotated, p-value of 0.173/0.002 in Wilcoxon test) (Figure 2d, h, Supplementary Figure 3).

A.3 mSTZ dataset

A.3.1 Identifying disease-associated genes in the mSTZ dataset.

We used annotation-trainability analysis to identify genes associated with the disease state of individual mSTZ β -cells within the single-cell pancreatic dataset [18], in the sense that their expression levels are correlated (positively-associated) or anti-correlated (negatively-associated) with the confidence score assigned to each cell’s disease annotation (Methods 2.3; Figure 1g). Among the five top-ranking genes that are positively-associated with a disease state, we identified known β -cells de-differentiation marker genes, including *Gpx3* (Glutathione peroxidase 3)[14], *Cd81*[16], *Aldh1a3*[7], and *Aqp4*[16], which have all been found to be upregulated in β -cells from mSTZ-induced diabetic models[16]; and mitochondrial respiratory chain complex I component *Ndufa1*, which is implicated in glucose metabolism regulation[10]. Among the five top-ranking genes that are negatively-associated with a disease state, we identified known β -cell activity markers *Trpm5*[7] and *Ins1*[7]; mitochondrial respiratory chain complex IV component *Cox6a2*, which has been found to be downregulated in diabetic islets[9]; *Fam151a*, which is also downregulated in diabetic pancreas[1]; and *Tmem215*, whose function and possible association with diabetes should be further investigated. The full list of disease-association gene scores is provided in Supplementary Table 2.

A.3.2 Evaluating treatment effectiveness and individual cell response

The islets dataset [18] includes six subgroups of mSTZ diabetic models, each treated with a different compound or combinations thereof: vehicle (control), Insulin, Insulin/GLP-1/Estrogen, GLP-1, Estrogen, and GLP-1/Estrogen[18]. We ranked all mSTZ cells according to their confidence and variability scores (Figure 5a), and analyzed the distribution of score ranks in each treatment subgroup (Figure 5g). In principle, the observed heterogeneity in disease progression among mSTZ cells (Figure 5a-f) could reflect treatment-independent heterogeneity, treatment-dependent heterogeneity, or both. To evaluate treatment-independent heterogeneity, we inspected the vehicle-treated subgroup, which serves as a negative control. While most cells in this subgroup were assigned high-ranking confidence scores and low-ranking variability scores, we identified a sizable subgroup with low-ranked confidence scores and high-ranked variability scores (Figure 5g, left panel), indicating lower congruence with a disease label. This subgroup is enriched for cells expressing healthy-like levels of *Ins1* (Figure 5a.g, left panel; Supplementary Figure 7a, left panel) and low levels of *Serpina7* (Supplementary Figure 7b,c, left panels). Thus, cellular heterogeneity in disease progression is at least partially independent of treatment.

We followed by evaluating treatment-dependent heterogeneity, that is, the effectiveness of the different treatments. For the two subgroups that were treated with insulin, either alone or in combination with GLP-1 and Estrogen, the distribution of the confidence and variability scores skewed notably toward lower confidence scores, along with high expression of the *Ins1* gene and low expression of the *Serpina7* gene; in contrast, for the subgroups treated with only GLP-1 or Estrogen, the distribution resembled the vehicle-treated control, and for the subgroup treated with GLP-1/Estrogen, it resembled an intermediate between the insulin-treated and the vehicle-treated cells (Figure 5g,h; Supplementary Figure 7a-c,n). Indeed, cells treated with Insulin were distributed nearly exclusively between the intermediate (20.9%) or healthy-like (77.3%) clusters in the trainability-aware graph; in contrast, cells treated with either GLP-1 or Estrogen were overrepresented in the disease cluster (67.2% and 73.8%, respectively), in proportions similar to the vehicle-treated cells (69.0%), with most remaining cells being assigned to the intermediate cluster (GLP-1 26.3%, Estrogen 21.9%, vehicle 26.5%); and cells treated with both GLP-1 and Estrogen were overrepresented in the intermediate state (40.8%) compared with the vehicle, but only a small fraction were assigned to the healthy-like cluster (23.4%) (Figure 5h; Supplementary Figure 7n). These findings align with a previous study, where insulin therapy was shown to restore β -cell function, leading to diabetes remission and improved insulin secretion, and where the combination of GLP-1 and estrogen was shown to be more effective than either treatment alone [15].

Upon scrutinizing the raw confidence and variability map, we identified a small subpopulation comprising 43 cells that are particularly hard to learn due to their low variability and confidence

levels (confidence < 0.15; variability < 0.15). Among these, 41 cells were subjected to Insulin or GLP-1/Estrogen/Insulin treatments (Figure 5i), and interestingly, they resembled control cells across all nine marker genes (Figure 5j). The overwhelming dominance of insulin-treated cells in this healthy-like subpopulation may imply that insulin treatment is exceptionally effective at restoring a nearly-healthy cell state, which cannot be explained by treatment-independent heterogeneity in mSTZ cells.

We conclude that annotation-trainability analysis revealed that treatments of mSTZ mice that include Insulin transform a substantial portion of their β -cells into a nearly-healthy state, in contrast with the treatments that do not include insulin. Moreover, annotation-trainability analysis can be used not only for inferring disease state of individual cells, but also as a sensitive tool for evaluating treatment effectiveness and individual cell responses to the same treatment.

A.4 Benchmarking

Erroneously annotated cells. To benchmark the identification of erroneous annotations, we compared the results of Annotatability to scReClassify[11], a state-of-the-art method for the identification of misclassified cells using the two different classifiers which were used in the original study: SVM and random forest, as well as the results obtained by an additional KNN classifier which we implemented based on the scReClassify framework. Specifically, we implemented scReClassify in Python, and compared the inferred confidence vector by Annotatability to the inferred probability of the cell being misclassified by scReClassify, which is the cumulative probability of the cell belonging to cell types other than the corresponding input annotation.

For the benchmarking analysis (using four spatial datasets, as described in Results 1.2.3), we created a semi-synthetic setting, where we randomly perturbed the annotation of varying fractions of cells (10%, 20%, 30%, 40%, and 50%) to a different label than their input annotation, chosen uniformly at random. We randomly subsampled each of the spatial datasets to 10,000 cells, with the exception of the Visium dataset which originally contained 2800 cells.

Ambiguous annotation identification. To identify cells in an ambiguous state, we created a semi-synthetic simulation using two scRNA-seq datasets: one sampled from PBMCs (Results 1.2.2) and the other from pancreatic islets (Results 1.2.5). We used four different cell types for the simulation: α , β , γ , and δ cells for the pancreatic islets data, and four randomly sampled cell types out of the ones described in Results 1.2.2 for the PBMC data. Next, we randomly sampled the percentage of cells to be placed in an ambiguous state from a Gaussian distribution with a mean of 0.5 and a standard deviation of 0.1. To generate an ambiguous state for cell i which is randomly sampled, we perturb its expression profile A_i by linearly interpolating it with an expression profile A_j of cell j (where j is sampled uniformly at random from all cells of a different label than i), to get a perturbed, ambiguous expression profile $\tilde{A}_i$: $\tilde{A}_i = A_i * (1 - \alpha) + \alpha * A_j$, where α is sampled from a uniform distribution between 0 and 1. Then, using the original unperturbed annotated dataset, we computed the absolute value of Spearman rank correlation between $(1 - \alpha)$ and the confidence assigned to each cell.

Trajectory inference. We compared the confidence scores inferred for the EMT (Results 1.2.4) and the mSTZ (Results 1.2.5) datasets to diffusion pseudotime trajectory inference (DPT) based on geodesic distance along a neighbors graph[20, 19]. Annotatability outperformed DPT for the task of trajectory inference for both biological settings (Results 1.2.4, 1.2.5).

A.5 Sensitivity analysis

We conducted a sensitivity analysis over a simulated dataset consisting of cells belonging to distinct clusters as well as cells in intermediate states. Specifically, we generated four distinct clusters, each consisting of 1250 observations (cells) with 1,000 variables (gene expression levels), using Scanpy (scanpy.datasets.blobs) [19]. The clusters are modeled as Gaussian mixtures, each with a mean sampled uniformly from $[-10, 10]$ and a standard deviation of 1. We clipped all negative values to zero, followed by standard preprocessing (Supplementary 2.7; Supplementary Figure 1a). To introduce intermediate states, we perturbed half of the cells in Cluster 1 by blending them with paired cells, all sampled uniformly at random, from Cluster 3. To blend the gene expression profile x_i^k of each perturbed cell i from cluster k with that of the gene expression profile x_j^l of its paired cell j from cluster l , we used linear interpolation as follows: $x_{new} = (1 - \alpha) * x_i^k + \alpha * x_j^l$, where x_{new} is the perturbed cell’s gene expression profile, and α is the mixing fraction (blending level), sampled uniformly from $[0, 1]$. This process effectively generated cells at intermediate states and mislabeled cells in the simulated dataset (as cells retained their original annotations), corresponding to medium and high mixing fractions, respectively (Supplementary Figure

1b,c). We then applied Annotatability to the blended dataset, with the following default parameters: learning rate = 0.001, batch size = 128, hidden layers = 2, epochs = 75. We inferred the confidence scores for all cells with respect to their input annotations, which resulted, as expected, with a lower confidence score with increasing mixture percentage (Supplementary Figure 1d). The results of Annotatability were robust to varying hyperparameter values; We observed consistently high Spearman’s rank correlation values between the inferred confidence scores and mixing fractions, ranging from 0.86 to 0.93 (Supplementary Figure 1), which were specifically robust to variations in learning rate (Supplementary Figure 1e), batch size (Supplementary Figure 1f)), number of hidden layers (Supplementary Figure 1g), and number of epochs (Supplementary Figure 1h).

A.6 Benchmarking Annotatability relative to AdaSampling

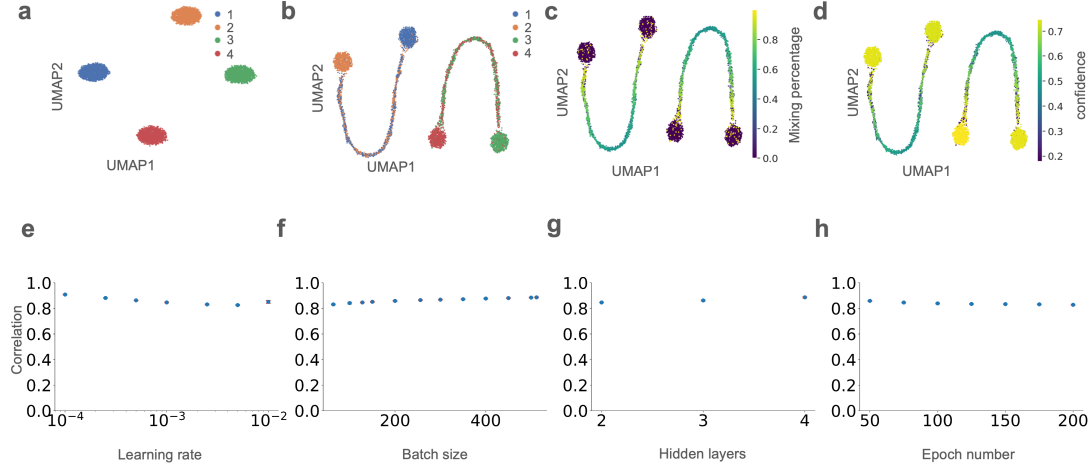
AdaSampling is a method that can be used to improve model training by reducing the impact of mislabeled data[11] and is tailored to re-annotation and detection of erroneous annotations. While for those tasks the performance of AdaSampling is similar to Annotatability (Figure 3d, Supplementary Figure 5e-g), Annotatability achieves much better results relative to AdaSampling in inferring trajectories of intermediate states between annotations (Supplementary Figure 4). In addition, Annotatability generates two metrics, confidence and variability, while AdaSampling, similarly to other probabilistic classifiers, provides only a single metric analogous to confidence. This difference is crucial since Annotatability’s two-dimensional metric allows to robustly categorize the data into three groups: correctly annotated, ambiguously annotated, and erroneously annotated. For example, using the variability metric we can more easily distinguish between hard-to-learn cells (erroneous annotations) which have low variability and low confidence to the ambiguous cells which have medium or high variability (e.g. Figure 2c). From a runtime complexity perspective, the application of AdaSampling, as described in [11], requires training the classifier 30 times, whereas Annotatability necessitates only a single training phase. Therefore, we recommend using Annotatability for identification of ambiguous states, for analysis of continuous dynamical processes annotated with discrete labels, and for detection of erroneous annotations (and subsequent re-annotations) in noisy scenarios, such as spatial transcriptomics data (as described in Results 1.2.3).

A.7 Supplementary Table

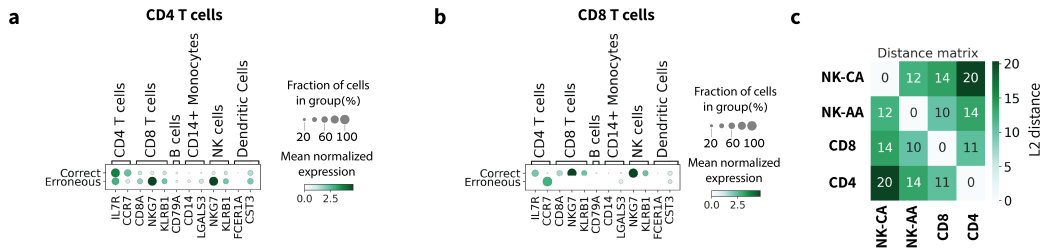
	Correctly annotated cells	Erroneously annotated cells	Re-annotated cells
Gad1-Inhibitory neurons	6.83	1.45	6.79
Slc17a6-Excitatory neurons	6.13	1.96	6.16
Myh11-Pericytes	7.27	-	4.77
Fn1-Endothelial	6.39	2.62	4.93
Cd24a- Ependymal	6.93	0	5.43
Selp1g- Microglia	7.84	0	5.89
Aqp4- Astrocytes	6.51	3.98	5.50
Pdgfra-OD Immature	7.11	0.55	4.76
Ttyh2- OD Mature	7.03	2.46	6.03
Mbp-OD Mature	3.84	1.73	3.25

Table 1: Mean expression of marker genes in correctly annotated cells, erroneously annotated cells and re-annotated cells in a MERFISH dataset of mouse hypothalamic preoptic region[13].

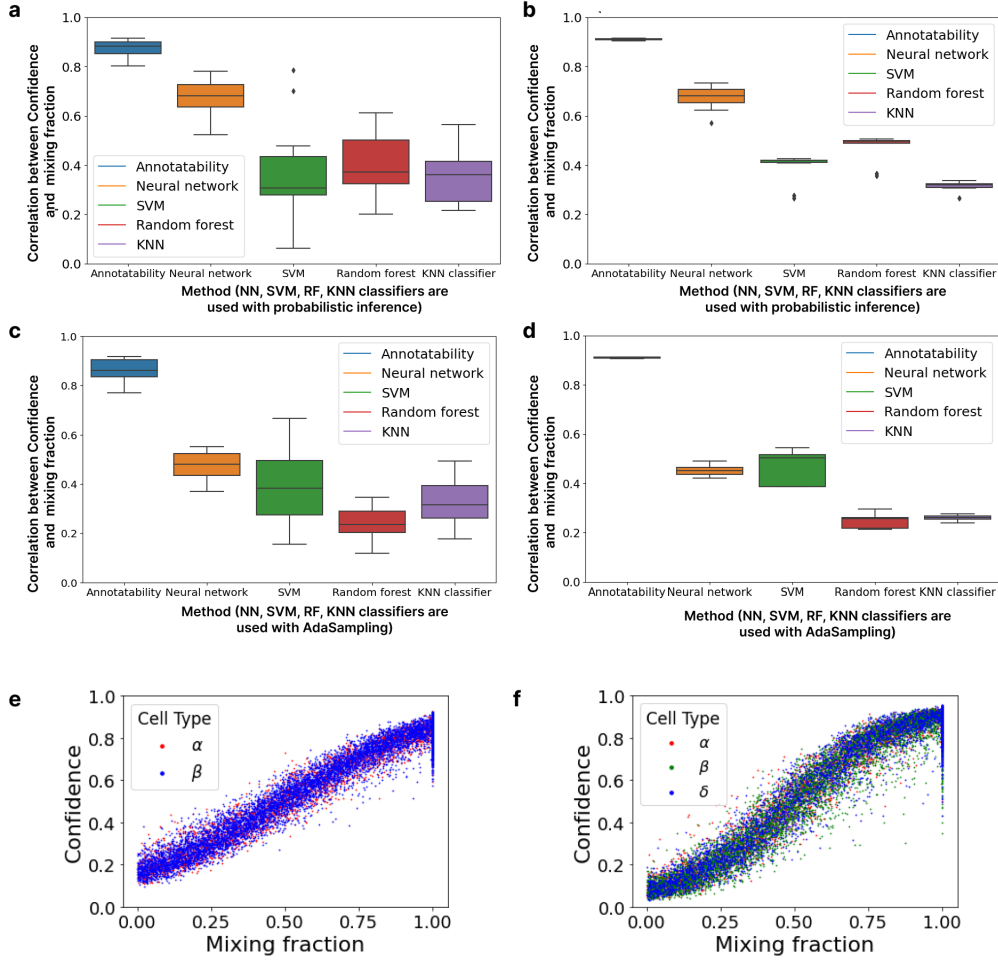
A.8 Supplementary Figures



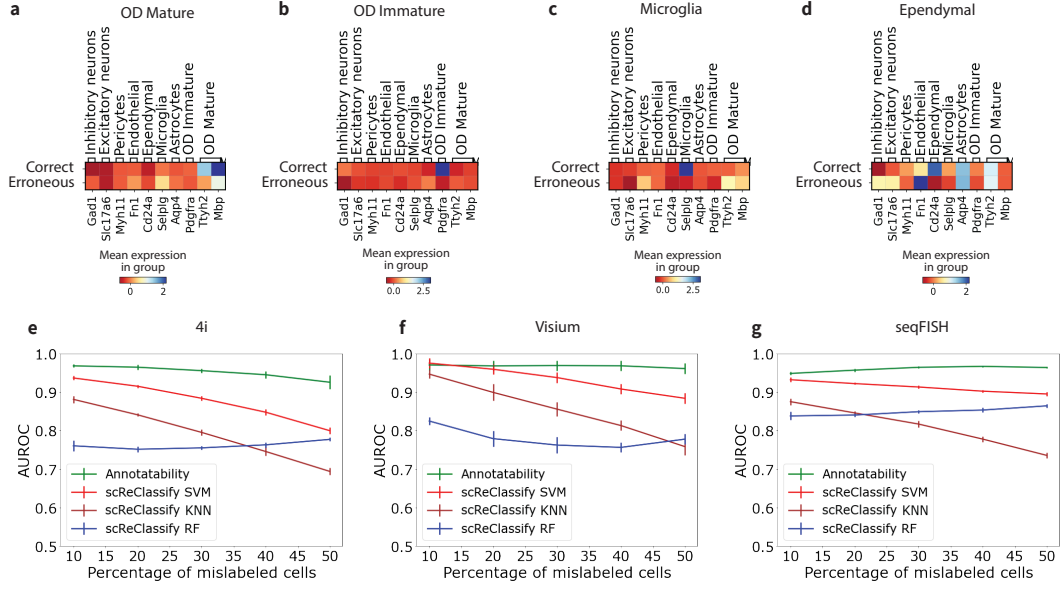
Supplementary Figure 1: Sensitivity analysis of Annotatability based on a simulated setting (Supplementary section A.5). (a) UMAP projections of the simulated data, comprising 1,000-dimensional gene expression profiles for 5,000 simulated cells, colored by the four assigned cell state annotations corresponding to each of the four cell clusters. (b) UMAP projections of the perturbed simulation data, where gene expression profiles for half of the cells were blended with those of random cells from neighboring clusters, using a randomly-selected mixing percentage (Supplementary section A.5). (c) UMAP, colored by the mixing percentage, indicating the fraction of the cell's expression levels originating from a cell in another cluster. (d) UMAP colored by confidence score. (e-h) The mean absolute Spearman's rank correlations between the confidence scores and the mixing percentages, presented as a function of various hyperparameters: learning rate (e), batch size (f), number of hidden layers (g), and number of epochs (h). For each set of hyperparameter values, the correlations are averaged over ten independent simulations. The default hyperparameters for the simulations are: learning rate = 0.001, batch size = 128, hidden layers = 2, epochs = 75. Error bars represent standard deviation between 10 different runs .



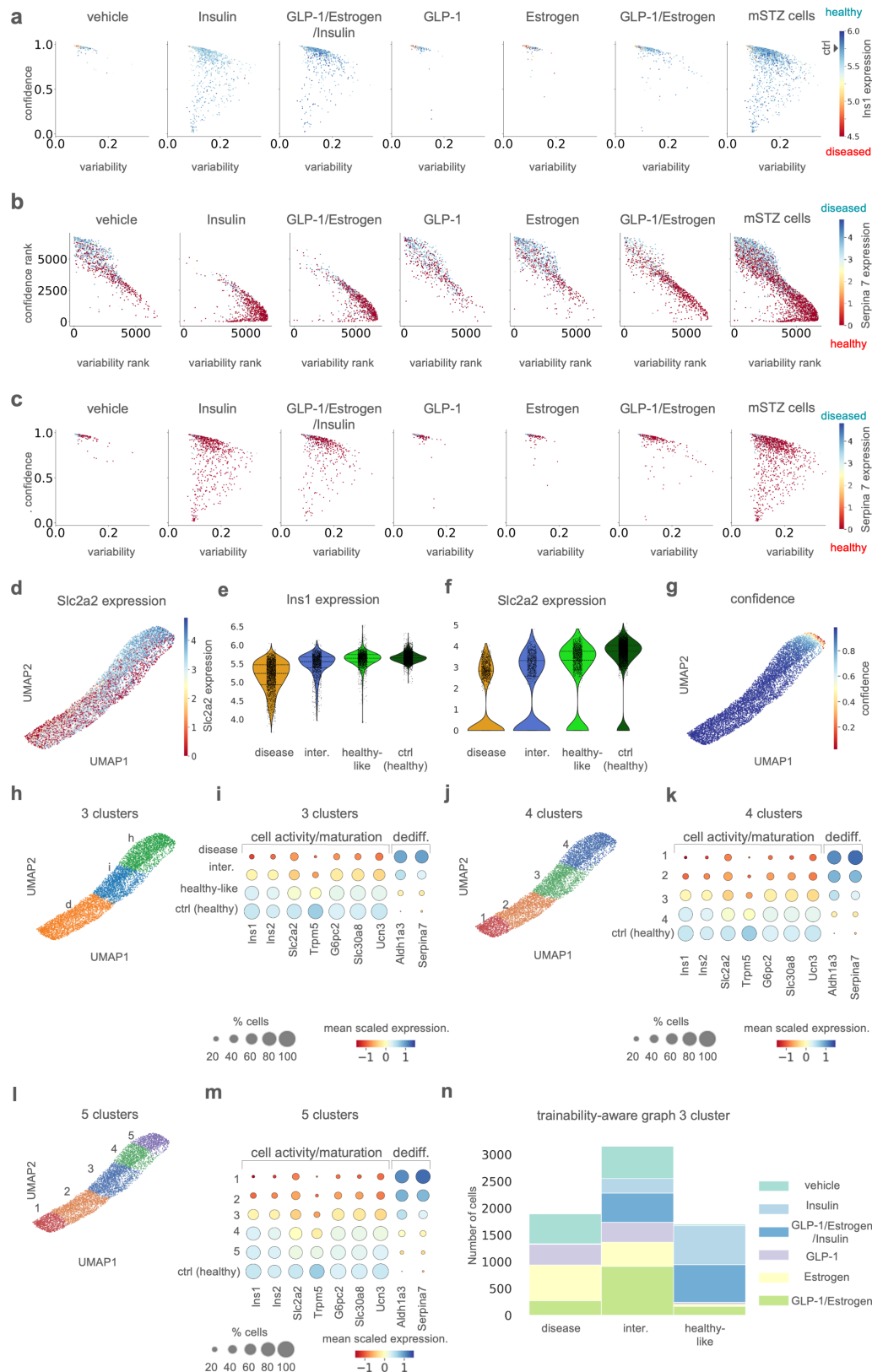
Supplementary Figure 2: (a, b) Dot plots of the expression of cell type marker genes of the annotated cell types in either CD4 T cells (a) or CD8 T cells (b) identified as correctly annotated (top row) or erroneously annotated (bottom row) in human PBMCs scRNA-seq data[4]. (c) L2 distance between the mean gene expression of the following four groups: NK cells which were classified as correctly annotated (NK-CA), NK cells which were classified as ambiguously annotated (NK-AA), CD8 T cells (CD8), and CD4 T cells (CD4).



Supplementary Figure 4: Benchmarking the identification of intermediate cell states based on a semi-synthetic setting based on the PBMC [5, 6, 4], and mSTZ datasets [18] (Supplementary A.4). (a-d) Spearman correlation between confidence score and mixing fraction (α), where confidence is inferred by Annotatability, and four additional classifiers: Neural network (with the same architecture and training procedure as Annotatability), KNN classifier, SVM classifier and Random forest classifier, either based on probabilistic inference (top row) or AdaSampling (bottom row). Box plots show the median (center line), the 25th and 75th percentiles (box edges), whiskers extending to 1.5x the IQR, and outliers as points. Minima and maxima within the whiskers are included (N=20). (e) 2D scatter plot of the confidence, as inferred by Annotatability, vs. original gene expression percentage ($1 - \alpha$) for the mSTZ dataset cell types (α and β cells). (f) same as e), for mixing of α , β and δ cells.

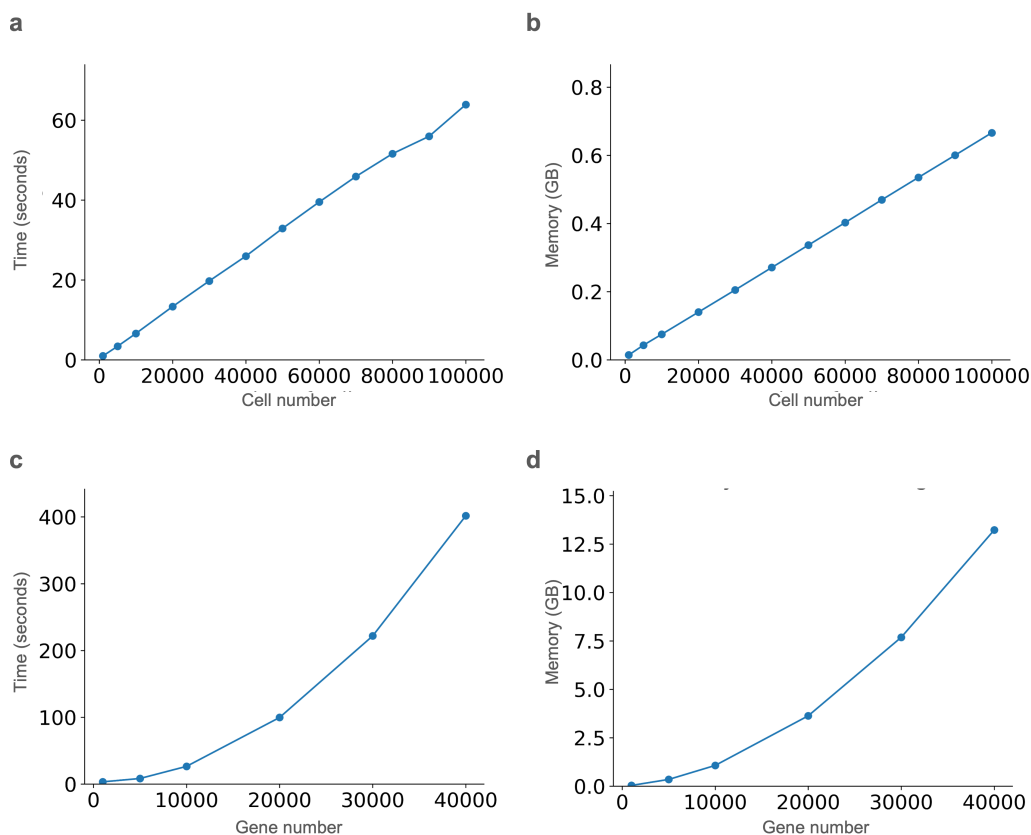


Supplementary Figure 5: (a-d) Heatmaps of the mean expression of marker genes of the annotated cell types, in OD Mature (a), OD Immature (b), Microglia (c), and Ependymal cells (d) identified as correctly annotated (top row) or erroneously annotated (bottom row) in MERFISH dataset of mouse hypothalamic preoptic region[13]. The cell type marker genes[13]: *Gad1* (Inhibitory neurons), *Slc17a6* (Excitatory neurons), *Myh11* (Pericytes), *Fn1* (Endothelial), *Cd24a*(Ependymal), *Selpg* (Microglia), *Aqp4* (Astrocytes), *Pdgfra* (OD Immature), *Ttyh2*, *MBbp* (OD Mature). (e) AUCROC vs percentage of mislabeled cells, for three spatial datasets: Four i (43 genes), Visium (16562 genes), and Seqfish (351 genes), as identified by four different methods: KNN classifier (brown), scReClassify with random forest classifier (blue), scReClassify with SVM classifier (red), and Annotatability(green). Dots represent mean value, vertical lines mark standard deviation.



Supplementary Figure 7: (a) 2D map of the confidence and variability scores with respect to a disease annotation for pancreatic islets scRNA-seq data[18], colored by *Ins1* expression, for cells from mSTZ diabetic models receiving different follow-up treatments (six panels starting from the left), and across all mSTZ cells regardless of follow-up treatment (rightmost panel).

Supplementary Figure 7: (b) 2D map of the ranks of the confidence and variability scores with respect to a disease annotation colored by *Serpina7* expression, for the same sets of cells as in (a). (c) Same as the 2D map in (a), colored by *Serpina7* expression. (d) 2D UMAP showing the trainability-aware graph for all mSTZ cells, colored by *Slc2a2* expression. (e) Violin plots showing the distribution of *Ins1* expression in cells from each of the three Louvain clusters obtained from the trainability-aware graph of the mSTZ cells, or from the healthy control. (f) Violin plots as in (e) for *Slc2a2* expression. (g) 2D UMAP as in (d), colored by confidence scores for all mSTZ cells. (h) 2D UMAP as in (d), colored by three Louvain clusters. (i) Dot plots showing the mean scaled expression of the same marker genes as in Figure 5c for the three Louvain clusters as in (h) vs. control healthy cells (bottom row). (j) Same 2D UMAP as in (h), shown for four clusters; the number of clusters was increased by tuning the Louvain clustering hyperparameters. (k) Same dot plots as in (i) for four Louvain clusters. (l) Same 2D UMAP as in (h), shown for five clusters. (m) Same dot plots as in (i) for five Louvain clusters. (n) Stacked bar plot showing the number of cells in each of the three Louvain clusters for each treatment.



Supplementary Figure 8: Scalability analysis of Annotatability based on a simulated setting (Supplementary A.5). (a-b) DNN training time and memory usage as a function of the number of cells, for a fixed number of genes ($n=1000$). (c-d) DNN training time and memory usage as a function of the number of genes, for a fixed number of cells ($m=5000$). All runs were performed on a single GPU, NVIDIA RTX A5000.

Supplementary References

- [1] Firoz Ahmed. “Deciphering the gene regulatory network associated with anti-apoptosis in the pancreatic islets of type 2 diabetes mice using computational approaches”. In: *AIMS Bioengineering* 10.2 (2023), pp. 111–140.
- [2] Megan A Cooper, Todd A Fehniger, and Michael A Caligiuri. “The biology of human natural killer-cell subsets”. In: *Trends in immunology* 22.11 (2001), pp. 633–640.
- [3] Aharon G Freud et al. “The broad spectrum of human natural killer cell diversity”. In: *Immunity* 47.5 (2017), pp. 820–833.

- [4] Adam Gayoso et al. “A Python library for probabilistic analysis of single-cell omics data”. In: *Nature biotechnology* 40.2 (2022), pp. 163–166.
- [5] 10x Genomics. *4k PBMCs from a Healthy Donor*. Version Cell Ranger 2.1.0. URL: <https://www.10xgenomics.com/datasets/4-k-pbm-cs-from-a-healthy-donor-2-standard-2-1-0>.
- [6] 10x Genomics. *8k PBMCs from a Healthy Donor*. Version Cell Ranger 2.1.0. URL: <https://www.10xgenomics.com/datasets/8-k-pbm-cs-from-a-healthy-donor-2-standard-2-1-0>.
- [7] Max Hahn et al. “Topologically selective islet vulnerability and self-sustained downregulation of markers for β -cell maturity in streptozotocin-induced diabetes”. In: *Communications Biology* 3.1 (2020), p. 541.
- [8] Yuhan Hao et al. “Integrated analysis of multimodal single-cell data”. In: *Cell* 184.13 (2021), pp. 3573–3587.
- [9] Elizabeth Haythorne et al. “Diabetes causes marked inhibition of mitochondrial metabolism in pancreatic β -cells”. In: *Nature communications* 10.1 (2019), p. 2474.
- [10] Wo-Lin Hou et al. “Inhibition of mitochondrial complex I improves glucose metabolism independently of AMPK activation”. In: *Journal of cellular and molecular medicine* 22.2 (2018), pp. 1316–1328.
- [11] Taiyun Kim et al. “scReClassify: post hoc cell type classification of single-cell rNA-seq data”. In: *BMC genomics* 20.9 (2019), pp. 1–10.
- [12] José L McFaline-Figueroa et al. “A pooled single-cell genetic screen identifies regulatory checkpoints in the continuum of the epithelial-to-mesenchymal transition”. In: *Nature genetics* 51.9 (2019), pp. 1389–1398.
- [13] Jeffrey R Moffitt et al. “Molecular, spatial, and functional single-cell profiling of the hypothalamic preoptic region”. In: *Science* 362.6416 (2018), eaau5324.
- [14] Lena Oppenländer et al. “Vertical sleeve gastrectomy triggers fast β -cell recovery upon overt diabetes”. In: *Molecular Metabolism* 54 (2021), p. 101330.
- [15] Stephan Sachs et al. “Targeted pharmacological therapy restores β -cell function for diabetes remission”. In: *Nature Metabolism* 2.2 (2020), pp. 192–209.
- [16] Ciro Salinno et al. “CD81 marks immature and dedifferentiated pancreatic β -cells”. In: *Molecular Metabolism* 49 (2021), p. 101188.
- [17] Daniel Schwabe et al. “The transcriptome dynamics of single cells during the cell cycle”. In: *Molecular systems biology* 16.11 (2020), e9946.
- [18] Sophie Tritschler et al. “A transcriptional cross species map of pancreatic islet cells”. In: *Molecular Metabolism* 66 (2022), p. 101595.
- [19] F Alexander Wolf, Philipp Angerer, and Fabian J Theis. “SCANPY: large-scale single-cell gene expression data analysis”. In: *Genome biology* 19 (2018), pp. 1–5.
- [20] F Alexander Wolf et al. “PAGA: graph abstraction reconciles clustering with trajectory inference through a topology preserving map of single cells”. In: *Genome biology* 20 (2019), pp. 1–9.
- [21] Rachel K Zwick et al. “Epithelial zonation along the mouse and human small intestine defines five discrete metabolic domains”. In: *Nature Cell Biology* (2024), pp. 1–13.