



Universiteit
Leiden

The Netherlands

Visual analytics for spatially resolved omics data at single cell resolution: methods & applications

Somarakis, A.

Citation

Somarakis, A. (2022, January 20). *Visual analytics for spatially resolved omics data at single cell resolution: methods & applications*. Retrieved from <https://hdl.handle.net/1887/3250550>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/3250550>

Note: To cite this publication please use the final published version (if applicable).

3

MICRO-ENVIRONMENT EXPLORATION

This chapter was adapted from:

A. Somarakis, V. V. Unen, F. Koning, B. P. Lelieveldt, and T. Höllt, “ImaCytE: Visual Exploration of Cellular Micro-Environments for Imaging Mass Cytometry Data,” *IEEE Transactions on Visualization and Computer Graphics*, vol. 27, no. 1, pp. 98–110, 2021.

3.1. ABSTRACT

Tissue functionality is determined by the characteristics of tissue-resident cells and their interactions within their microenvironment. Imaging Mass Cytometry offers the opportunity to distinguish cell types with high precision and link them to their spatial location in intact tissues at sub-cellular resolution. This technology produces large amounts of spatially-resolved high-dimensional data, which constitutes a serious challenge for the data analysis. We present an interactive visual analysis workflow for the end-to-end analysis of Imaging Mass Cytometry data that was developed in close collaboration with domain expert partners. We implemented the presented workflow in an interactive visual analysis tool; ImaCytE. Our workflow is designed to allow the user to discriminate cell types according to their protein expression profiles and analyze their cellular microenvironments, aiding in the formulation or verification of hypotheses on tissue architecture and function. Finally, we show the effectiveness of our workflow and ImaCytE through a case study performed by a collaborating specialist.

3.2. INTRODUCTION

Cells are the structural units of life and the main orchestrators of tissue function [1]. In recent years it has become clear that the phenotype and function of cells is co-determined by their location and interactions within the tissue context. Detailed analysis of the heterogeneity of cells provides information on tissue composition, while the spatial organization of cells provides information on tissue function. Both aspects are important for unraveling the complexity of tissue function and disease.

Conventional analysis of cell heterogeneity requires processing of tissue specimens into single-cell suspensions at the expense of spatial information. Traditional imaging analysis based on immunofluorescence is limited by the number of proteins that can be analyzed simultaneously and consequently the spatial location of only a few cell types can be revealed. A novel imaging modality, Imaging Mass Cytometry [2] was recently introduced allowing the measurement of the expression of up to 40 proteins simultaneously with a spatial resolution as low as one micrometer per pixel, preserving the tissue architecture at sub-cellular resolution [3]. The broad range of measured proteins allows the identification of a large variety of cell phenotypes within the tissue context. However, the wide range of discovered cell phenotypes in combination with the spatial resolution creates a highly complex system to analyze.

For analysis, each measured protein is typically interpreted as a dimension, defining every cell as a data-point in a high-dimensional space. The analysis of this high-dimensional space allows for the identification of distinct cell phenotypes, not known *a priori*, similar to non-spatial methods. In addition, the spatial resolution of the data allows biologists to localize the identified cells in the tissue architecture and form hypotheses based on the location of cells and their microenvironment. Here, we define the microenvironment of a cell as the directly adjacent cells in the tissue.

To explore the cells in their corresponding microenvironments, we first need to identify the different cell phenotypes existing in the analyzed tissue. Therefore, we extended our previous work [4] focused on the phenotype identification of non-spatial mass cytometry data to enable the exploration of cellular microenvironments and discover subsets of cell

phenotypes with unique microenvironment characteristics. We propose an interactive, data-driven workflow and consequently a tool that has been designed for the end-to-end exploratory analysis of Imaging Mass Cytometry data. The main contributions of this paper are:

- 1 An end-to-end workflow for the analysis of spatially resolved *-omics* data, including
 - interactive definition of cohesive phenotypical groups,
 - stratification of phenotypical groups based on their microenvironment characteristics,
 - inspection of significant spatial interactions in the tissue, and
 - motifs to group and visualize cells with a similar cellular microenvironment.
- 2 The implementation of the proposed workflow in an interactive visual analysis tool, we call ImaCytE.

The remainder of this paper is structured as follows. We present related work in Section 3.3, followed by a brief description of the biological background in Section 3.4. In Section 3.5 we identify and abstract tasks and present their corresponding designs. The effectiveness of our tool and workflow is shown in a case study in Section 3.6. We conclude in Section 3.7 and present open problems and potential directions for future work.

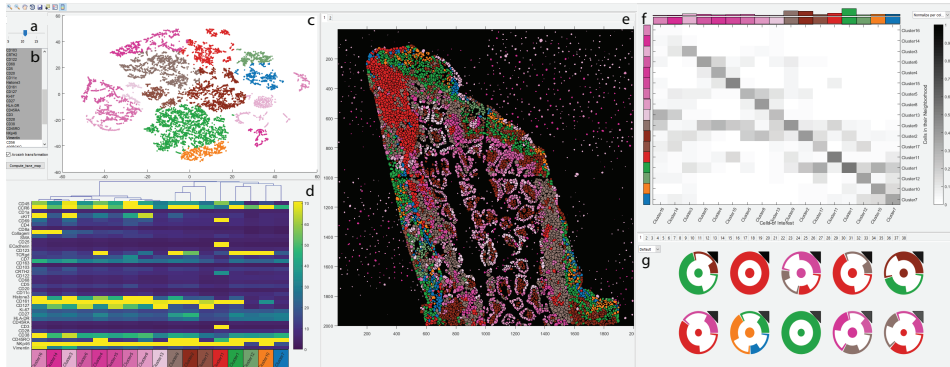


Figure 1: ImaCytE. Screenshot of our integrated system including the settings panels (a,b), the embedding view (c), the heatmap view (d) for cluster visualization, the image view (e), the interactions heatmap (f), and the motif view (g).

3.3. RELATED WORK

The routine acquisition of high-resolution spatially-resolved, *-omics* data is a relatively new development. Consequently, only few integrated analysis tools, leveraging the full complexity of such data exist. Multeesum [5] allows the comparative visualization of

spatio-temporal gene expression in fruit fly embryos acquired through photon microscopy. Abdemoula et al. [6], used a t-SNE-based workflow for the exploration of Imaging Mass Spectrometry data [7]. InsituNet [8] is using an interactive network-based visualization to illustrate the spatial co-expression of transcripts over a specified region of the tissue. Even though InsituNet allows the location of transcript interactions, one cannot explore spatial co-expression of transcript subsets. Compared to Imaging Mass Cytometry, all of these techniques deal with relatively low resolution data and none of them provides means to interpret microenvironment patterns.

Phenotype identification in non-spatially resolved Mass Cytometry [9] data has been a very active field of research, resulting in a large amount of tools. The most widely used ones include Vortex [10], FlowMaps [11], FlowSOM [12], and Phenograph [13], all employing unsupervised clustering. Amir et al. proposed viSNE [14], a manual technique for cell phenotype identification based on dimensionality reduction using t-SNE [15], through which local structure is visible at single-cell resolution. However, manual identification of phenotypically distinct clusters of cells is cumbersome and time-consuming. In previous work [4, 16], we developed Cytosplore, following the progressive visual analytics paradigm [17, 18], to allow interactive phenotype specification and adjustments. Here, we extend the concepts presented in the original work on Cytosplore to support the analysis of Imaging Mass Cytometry data.

To the best of our knowledge, the only integrated tool for phenotype identification and microenvironment analysis for Imaging Mass Cytometry data is histoCAT [19]. histoCAT provides a relatively fixed analysis pipeline, based on automated clustering for the phenotype identification and statistical analysis to identify samples with significant interactions among phenotypes. While histoCAT provides visualizations to present the results of the steps of the pipeline, it does not allow for fully interactive exploration, as we propose here. In particular, a key differentiator of our proposed workflow is that we link abstract information, such as protein expression or the identified phenotype, to the spatial information throughout all steps of the analysis. This makes an interactive, exploratory workflow possible that allows the user to make informed decisions at every step, including quality control, phenotype identification, and microenvironment characterization. This interactive workflow also allows for one of the main contributions of this work, the exploration of cell microenvironments and interactions. While histoCAT only allows classification of complete samples, based on automatically identified interaction patterns, here, we propose a glyph-based visual encoding to make the exploration of different patterns possible, link them to their location in tissue and finally stratify cell phenotypes further, based on their microenvironments.

To aggregate and compare different cell microenvironments we group them into motifs which we visualize with a glyph-based [20, 21] representation. Similar interaction patterns, such as protein interaction networks are often visualized with node-link diagrams. Such visualizations are insightful when the number of edges is comparable to the number of nodes or a pathway continuity should be depicted. Otherwise, the outcome is a complex and crowded network [22]. Our glyph representation for the cell microenvironment networks is inspired by Landesberger et al. [23] and Dunee et al. [24], who use motifs in order to simplify and enhance their network visualizations and for illustrating the network variance, respectively. However, these visualizations are mostly focused on the way that the data are

connected (i.e., star, clique etc.) and not on the type of the data (i.e. clusters), as is the case for the data presented in this work.

3.4. BACKGROUND

In recent years, high resolution and highly multiplexed imaging techniques have become available. State-of-the-art techniques like FISSEQ [25], smFISH [26], Padlock probes and RCA [27] allow transcriptome measurements at sub-cellular resolution. In combination with appropriate data analysis techniques these data are currently revolutionizing our perception of complex biological systems. Similarly, the newly introduced Imaging Mass Cytometry [2] produces multiplexed measurements of protein abundance at sub-cellular resolution.

3.4.1. DATA ACQUISITION

Imaging Mass Cytometry data acquisition consists of three steps. First, tissue sections are stained with antibodies conjugated to heavy metals that bind to specific proteins expressed by cells or the extracellular matrix. Second, the tissue sections are ablated spot by spot in a regular grid, where each spot represents a pixel in the resulting image. Currently the resolution of a pixel is $1\mu m$. Finally, the ablated material is guided to a traditional Mass Cytometer [9], where the abundance of each metal and thereby the expression of the corresponding protein are measured per pixel. Up to now, 40 different proteins can be simultaneously measured, but this number could increase up to 100 when appropriate metal reagents become available. Ultimately, the output of this process is a stack of gray-scale images, each depicting the abundance of one measured protein. While in principle, the values per pixel and protein correspond to the number of measured heavy metal ions, due to measurement limitations the actual values are continuous, to compensate for spillover between pixels and similar effects.

3.4.2. DATA PREPROCESSING

Although the output of the Imaging Mass Cytometry is a stack of gray-scale images with per-pixel measurements, ultimately the contained cells are of interest. In order to extract single-cell information from these images these need to be segmented in a preprocessing step. The segmentation process consists of two parts, as proposed by Shapiro et al. [19]. First, Ilastik [28], a semi-supervised machine learning technique is used to create a probabilistic classification per pixel. Subsequently, these classification masks are loaded into CellProfiler [29] which derives a segmentation mask. Based on the extracted segmentation mask, the protein abundance for each segmented cell can then be aggregated. The result is a simple table that contains the expression of all proteins per cell and the mask that relates every cell to a set of pixels in the original image. Based on the segmentation mask, we define the microenvironment of a segmented cell as the directly adjacent cells in image space.

3.5. IMACYTE

We designed and implemented our interactive visual analysis tool ImaCytE to support three main tasks;

T1 Quality Control,

T2 Cell Phenotype Identification, and

T3 Cell Microenvironment Exploration.

We follow Brehmer and Munzners task typology [30] for multi-level tasks and the extension for tasks regarding the high-dimensional data analysis [31]. In the following, we use a mono-spaced font when referencing this typology.

An overview of the three identified tasks is illustrated in Figure 2. During quality control (Figure 2a, Section 3.5.1) uninformative samples and proteins are identified and excluded from further analysis. Afterwards, the phenotype of each cell is identified (Figure 2b, Section 3.5.2), and finally the spatial interactions among the cells in the microenvironment are explored (Figure 2c, Section 3.5.3) by the user. In the following we discuss each of the tasks, and sub-tasks where available, in detail.

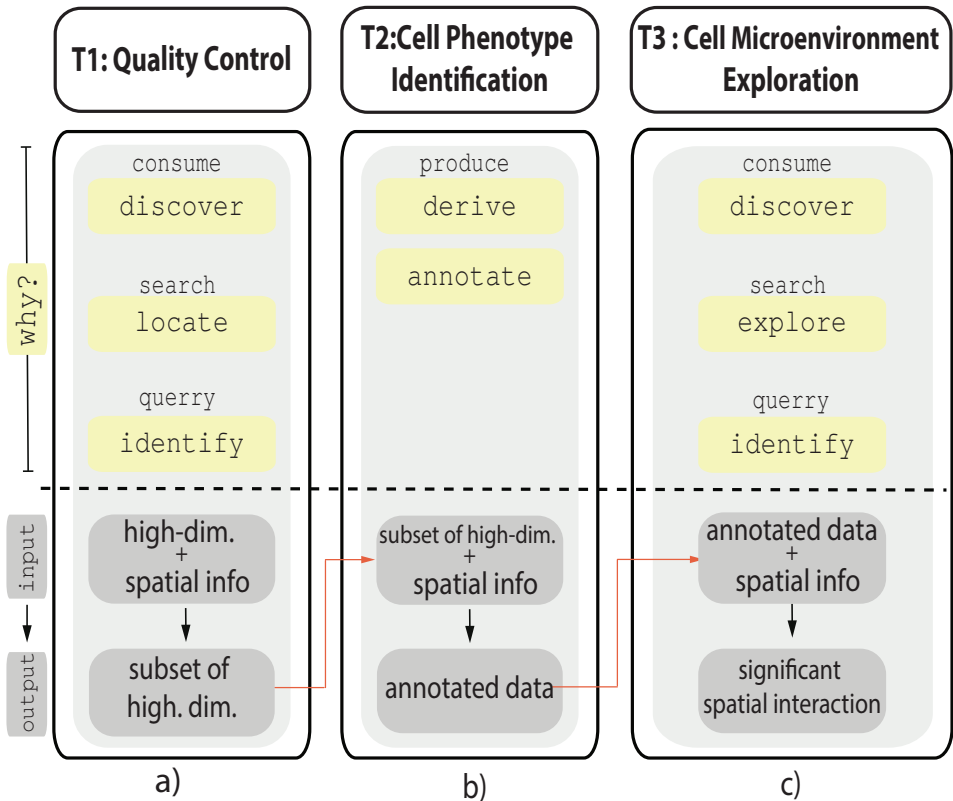


Figure 2: High-Level Task Overview. (a) T1:Quality Control, (b) T2: Cell Phenotype Identification and (c) T3: Cell Microenvironment Exploration.

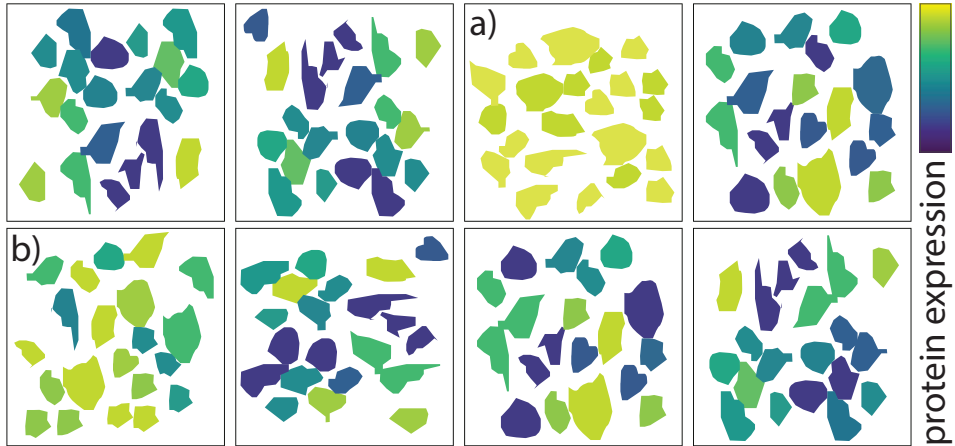


Figure 3: Batch Effects visible in small multiples view. a) Overly expressed marker (all cells indicate close to maximum expression). b) Offset caused by background staining (no low expression cells).

3.5.1. T1: QUALITY CONTROL

Imaging Mass Cytometry data are acquired in patches where an edge typically measures approximately one millimeter. Larger tissue regions can be combined from multiple patches, or samples from different regions are combined for cohort analysis. Combining multiple samples in this fashion makes the data prone to batch effects, caused by differences in handling or preservation of the tissue and are indicated by variation in the staining efficiency between samples as illustrated in Figure 3. Typical examples are background staining [32] causing an offset in protein expression (Figure 3b) or overly abundant expression of a protein caused by universal binding of the protein (Figure 3a).

Typically, protein markers are validated separately by conventional immunocytochemistry before being used in Imaging Mass Cytometry, to make sure that expression patterns match those reported in literature. However, not all markers work equally well in the staining process, leaving uncertainty whether they are functional in Imaging Mass Cytometry. In order to evaluate their functionality, the user must be able to verify each marker and its expression patterns for all samples.

Abstraction. During Quality Control (Figure 2a) the user *discovers*, *locates*, and *identifies* samples as well as protein markers that need to be excluded from the analysis. The input is the complete dataset as described in Section 3.4.2 consisting of the cells and their high-dimensional marker expression in combination with their location in the tissue. The output is a set of samples and markers to be used for further analysis.

ImaCytE. The core of quality control is the visualization of marker expression for all cells in their spatial context. In the most simple case the user can select a sample alongside a protein marker to visualize the expression of that marker in the spatial context. We use the original cell segmentation mask and use color-coding to encode the values of a selected protein expression for a selected sample as more effective channels such as position [33] are already blocked by the spatial nature of the data. By default, we use the viridis [34] colormap since it provides a perceptually uniform representation of the data, minimizes

contrast artifacts, and is color blindness friendly. Inspecting one sample and one marker at a time to identify differences between samples/markers heavily relies on user memory and can be challenging. Therefore, we provide the user with the possibility to additionally select either the expression of a single protein for one or more samples (`discover batch effects`) or the expression of one or more proteins for a specific sample (`discover broken markers`) in a small multiples [35] view as illustrated in Figure 3. However, selection of multiple proteins alongside multiple samples is not possible.

In a typical exploration session, the user first selects a protein of interest and all samples for display in the small multiples image viewer. Then, she `locates` and `identifies` potential samples where the staining process was not effective. Such samples can then either be excluded from further analysis or adjusted in external tools. To verify working markers, the user selects a sample and displays the different markers for the same sample. In case a non-functional marker can be identified in this process it can be ignored for further computational steps, such as clustering or dimensionality reduction.

3.5.2. T2: CELL PHENOTYPE IDENTIFICATION

Task T2, illustrated in Figure 4, describes the identification and labeling of each cell with its corresponding phenotype. Here, the phenotype is defined by the expression of different proteins by the cell similar to conventional Mass Cytometry. Therefore, we adapt the phenotype identification process described in our earlier work [4] to the specific needs of Imaging Mass Cytometry. In particular, since we generally observe significantly fewer cells in Imaging Mass Cytometry (tens of thousands compared to millions in conventional Mass Cytometry), we can skip the lineage delineation step [4, Task 1]. Instead, we extend the phenotype identification step [4, Task 2] by a semantic grouping of subsets and the possibility to inspect the resulting phenotype definition directly in the tissue images.

As illustrated in Figure 4, T2 is further divided into three sub-tasks. The first step, T2.a, is dimensionality reduction of the data for visualization. In the second step, T2.b, we then cluster the dimensionality reduced data to define groups of phenotypically similar cells. Finally, in step T2.c, the clusters are verified, labeled, and semantically grouped. In the following, we describe each sub-task separately. Since tasks T2.a and T2.b are identical to our previous work we provide a brief description only.

3.5.2.1. T2.A: DIMENSIONALITY REDUCTION

Abstraction. Here, we `derive` a two-dimensional embedding of the cells for visualization. The input are the cells (data-points) originating from the samples selected in task T1 alongside the markers (high-dimensional space), identified to be functional in T1.

ImaCytE. The goal of the dimensionality reduction is to provide the user with a low-dimensional representation of the data that allows the visual identification of phenotypically similar cells. As described in our previous work [4], t-SNE-based approaches [15] are a good choice for this task, as they aim to preserve neighborhoods of the original space in the low-dimensional space. Linderman and Steinerberger show that t-SNE preserves well separated clusters in high-dimensional data and relates to spectral clustering in the high-dimensional space [36]. This means that cells that are similar according to their high-dimensional protein expression will be close together in the low-dimensional embedding, making such groups easy to identify, for example in a two-dimensional scatter plot. We use

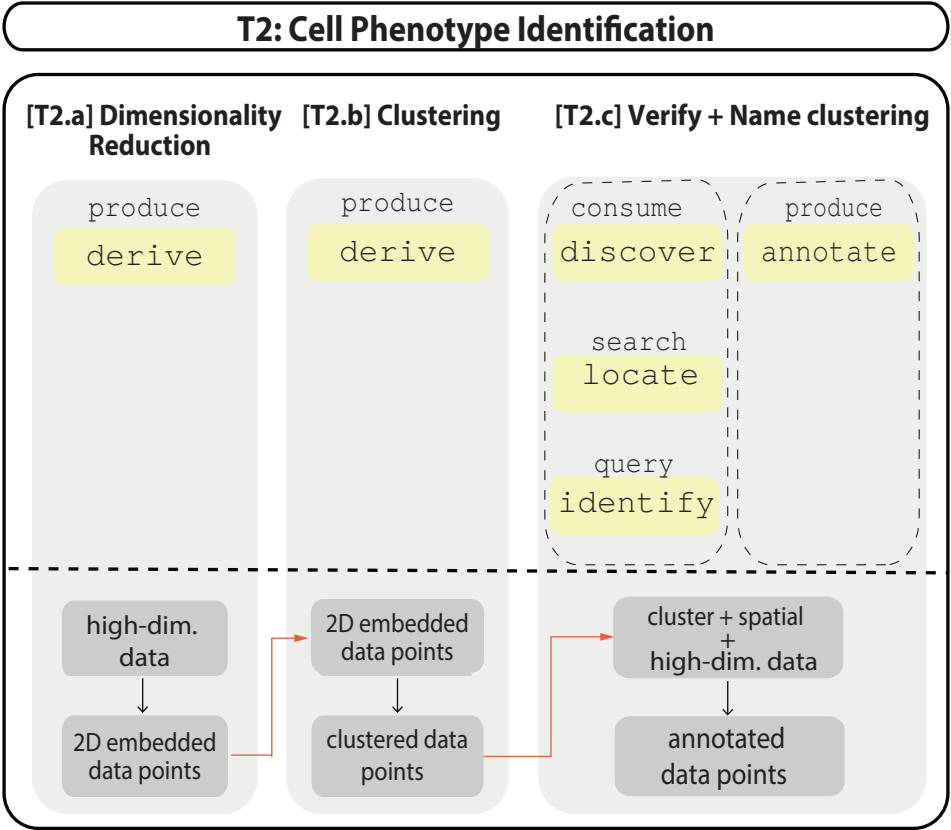


Figure 4: *The Cell Phenotype Identification Task* is divided into three sub-tasks, dimensionality reduction, clustering, and verification and annotation.

the A-tSNE algorithm [37] to `derive` such an embedding, as it provides highly similar results as the original t-SNE in a fraction of the time, making it viable for interactive exploration.

3.5.2.2. T2.B: CLUSTERING

Abstraction. In addition to the visualization of the cells according to their similarities we also want to `derive` clusters of similar cells for quicker labeling in the following step. Here, we use the dimensionality-reduced data created in the previous step as input, to make sure derived clusters match the visually identified groups in the previous step. However, in principle other forms of clustering using the complete high-dimensional information are possible. Those would require the high-dimensional data as input. The output is the clustered data, i.e., each cell annotated with a cluster id.

ImaCytE. In our implementation, we use the Mean-Shift algorithm [38] for clustering the dimensionality reduced data. Mean-Shift clustering is density-based and therefore the resulting logical clusters closely resemble visual clusters in the embedding. Furthermore, it can extract arbitrarily shaped clusters and does not require the specification of the number

of clusters in advance. Given an existing two-dimensional embedding, it can be computed very quickly. All these properties make it a suitable technique for defining clusters in this setting. It should be noted, however, that there are potential issues the user needs to be aware of. In brief, the two-dimensional embedding might not be able to sufficiently extract all meaningful structure of the high-dimensional space, resulting in heterogeneous visual clusters. Furthermore, particularly for very large data, the optimization process of t-SNE tends to end in local minima, 'tearing' high-dimensional clusters into two or more visual clusters. For a more detailed discussion, we refer to our previous work [4]. As indicated above, in principle any kind of clustering could be employed in this step and combined with the visualization in the two-dimensional embedding and in our implementation, we allow loading of clusterings created outside of our application to replace this step if desired. In practice, however, we [4, 16] and others [14, 39] have found that clustering t-SNE embeddings of Mass Cytometry data produces biologically meaningful results and in case structure is missed our interactive workflow allows for adjustments, such as merging of similar clusters in the next step.

3.5.2.3. T2.c: VERIFY AND NAME CLUSTERING

Abstraction. Task T2.c combines cluster verification and annotation. Based on the complete aggregated protein expression profile the user can `discover`, `verify` and `annotate` the derived clusters. Compared to its counterpart in our previous work, we add the possibility to `verify` the clustering in relation to the tissue and introduce a semantic, two-level `annotation` scheme.

The input are the clustered data-points from the previous step alongside their spatial as well as high-dimensional feature-space information. The output is a semantic label for each cell.

ImaCytE. The basis for the implementation of task T2.c again forms our implementation in Cytosplore [4, Task 2c]. As in our previous work, the expression of one protein can be used to color-code the embedding allowing the user to assess the homogeneity of the clusters with respect to the given protein. Furthermore, we use a cluster heatmap, showing the median expression of each protein per cluster allowing the biological interpretation of each cluster. Through the heatmap, the user can also directly adjust the clustering. When multiple clusters with a similar biological interpretation have been extracted they can be merged. This process is supported by sorting the clusters in the heatmap by similarity. Therefore, we compute a hierarchical clustering, based on the median expression values of each cluster and add a dendrogram to the heatmap.

We extend our original work in two ways. We introduce a semantic, two-level color-coding for the identified clusters, separating main cell types and sub-types. Making use of the spatial nature of the data, we extend our previous work with the possibility to verify all steps by visualizing the cluster result in the image view.

Since the number of cells in typical Imaging Mass Cytometry datasets is much smaller than in conventional Mass Cytometry, we do not need to divide the data as in our previous work [4, Task 1]. To retain the semantic separation of the data into main types, i.e., cell lineages, and sub-types of cells, we create a two-level color scheme illustrated in Figure 5. Based on the median expression of the derived clusters we compute a hierarchical clustering. Afterwards, a user-defined dissimilarity threshold (Figure 5b), based on the Euclidean distance in the cluster hierarchy, is used to separate the hierarchy into the main

cell types (clusters with large Euclidean distance, branching above the threshold) and sub-types (clusters with small Euclidean distance, branching below the threshold). We then assign semantic color-coding to the clusters according to their separation. We use the D3 [40] categorical colormap to assign primary colors to the main partitions (Figure 5c, red, blue, green) and vary the saturation of those colors within a partition (Figure 5c, shades of blue) to demonstrate the correlation of those clusters. Finally, we provide the option for the user to manually assign specific colors to the major groups to comply with the phenotypes and their biological significance.

As described above, the second addition is the introduction of the spatial information for verification of the resulting clusters. Here, in particular, we color the cells in the image view according to the color-scheme described above to verify the biological significance of the derived clusters. For example, if staining issues were missed in the quality control step it can happen that several clusters with the same biological meaning are created. Figure 6 illustrates such a case. Clusters C1 and C2 show similar protein expression patterns, with generally larger values for C2. Both clusters are visible in sample 1, indicating that they might represent the same phenotype with a large dynamic range. Sample 2 only exhibits cells from cluster C1, indicating that staining efficiency in that sample was worse than for sample 1 resulting in smaller dynamic range. With this information, the user can now decide to either go back to the quality control step and further investigate sample 2 or simply merge the two clusters to indicate their biological similarity.

All steps of Task T2 are completely interactive and linked, allowing iterative refinement of the clustering. Changes in clustering automatically trigger updates of the color-coding of the tissue in the spatial view, the embedding scatter plot and the heatmap. Thereby, the user can verify clustering results immediately across several linked views. Furthermore, the user can interactively probe in the heatmap, embedding, or spatial view through linked selections. Selecting one or multiple cells or clusters in any of the three views highlights the cells in any of the other visualizations. Highlighting in the heatmap and image views is achieved by fading out the non-selected area. In the scatter plot we change the mark from circles to crosses for the selection.

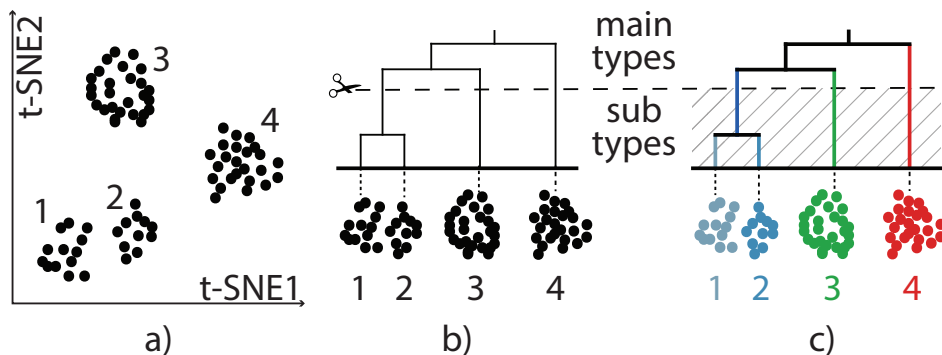


Figure 5: Two-Level Color-Coding. Clusters from the embedding, a), are hierarchically clustered, b), and a user-defined threshold is used to define main- and sub-types. Finally, color hues are assigned to the main-types and different saturation is applied to the corresponding sub-types, c).

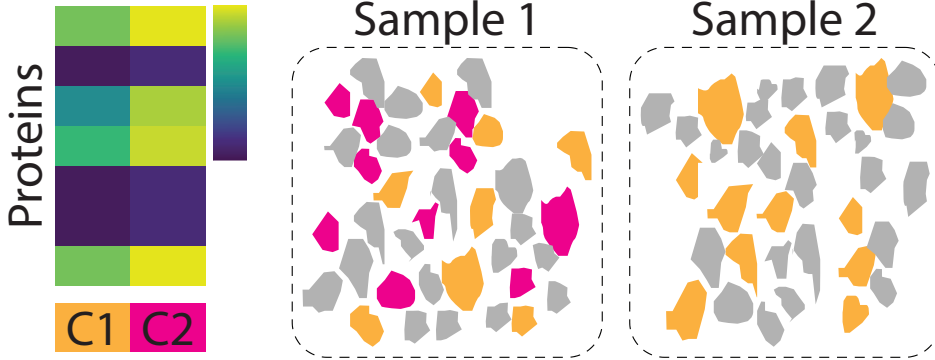


Figure 6: Variation in Staining Efficiency becoming apparent during cluster verification. Two clusters, C1 and C2, are identified, however, their protein expression, visualized in the heatmap, indicates an offset that can be caused by insufficient staining. Visualizing the clusters in their spatial context (grey indicates other cells) reveals that one of the clusters is indeed only existent in one of the samples and is likely caused by insufficient staining.

3.5.3. T3: CELL MICROENVIRONMENT EXPLORATION

Task T3 (Figure 7) is made possible due to the spatial resolution of the data and comprises the exploration of the cell microenvironment. In an overview first, detail on demand fashion, the task is divided into two sub-tasks. The first sub-task, T3.a Spatial Interactions Overview, is to gain an overview of which phenotypes interact with each other on a global scale. These interactions characterize the tissue function and can be helpful in order to categorize the tissue samples. However, a phenotypical subset can include thousands of cells with distinct local cell microenvironments, meaning subsets within a phenotype that have different spatial interactions from those illustrated in the overview are likely to exist. The differentiation of those subsets is of major interest to identify functionality of cells, beyond their phenotype. Sub-task T3.b, Spatial Interactions Details, describes the exploration of details of the interaction patterns on a local microenvironment scale.

For both sub-tasks it is important to link the abstract results of spatial interactions to the spatial location of the cells in the tissue. As an example our collaborators described *floating cancer cells*. Floating cancer cells are cancer cells that abandoned a tumor. These cells are prone to attacks from immune cells and as a result exhibit diverse microenvironments. Finding cells with a similarly diverse microenvironment within the tumor would have completely different biological implications. Therefore, it is important to provide the means to inspect the abstract results linked to tissue location. In summary, this example illustrates that we need to identify subsets of a phenotype with unique microenvironment characteristics, connect them to their spatial location and examine their characteristics separately, as they can be useful biomarkers for tissue functionality.

3.5.3.1. T3.A: SPATIAL INTERACTIONS OVERVIEW

Abstraction. For sub-task T3.a, the user *discovers* cells of which phenotypes interact spatially, *explores* how much they interact and *identifies* which of the interactions are significant. The input of this sub-task are all cells alongside their location in the tissue and their phenotype identified in Task T2. The output is the quantified amount of pairwise

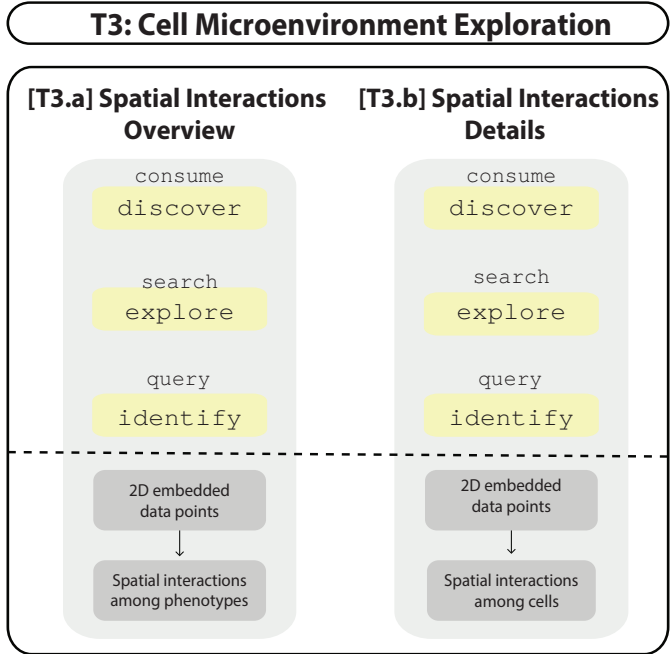


Figure 7: The Cell Microenvironment Exploration Task, including two sub-tasks. The first generates an overview of the spatial interactions among phenotypes. Afterwards, spatial interaction details on the cellular level are explored.

spatial interaction between phenotypes.

ImaCytE. Here, the goal is to provide a global overview of which phenotypes can be found in the microenvironment of cells with a specific phenotype. Therefore, for every cell we count the number of occurrences of each phenotype in its direct microenvironment and aggregate these counts over all cells of the same phenotype. The result is a directed and weighted graph, where each node represents a phenotype. The weighted link between two nodes illustrates the number of times a cell of the phenotype at the end of the link occurs in the microenvironment of a cell of the phenotype at the source of the link. In our current studies, the number of nodes n is typically in the range of 20 – 40. While pruning could be used to remove links with low weights, typically all nodes are connected to all other nodes, resulting in up to n^2 links. While intuitive to read, a node link diagram would be too dense and suffer in terms of readability [41]. Therefore, we chose a heatmap (Figure 8) to visualize the resulting graph, with the source nodes on the horizontal axis and the target nodes on the vertical axis. The weight or number of interactions is then indicated using color in the cell at the intersection of the source and target node. To not interfere with the qualitative colormap used to identify the different phenotypes, we use a white-to-black colormap to indicate the weight. The values in the heatmap are normalized. The user can choose different normalization criteria for the heatmap such as per row or per column. If the user is interested firstly in the different phenotypes that are most likely to exist in the microenvironment of a specific source cell normalization per column is

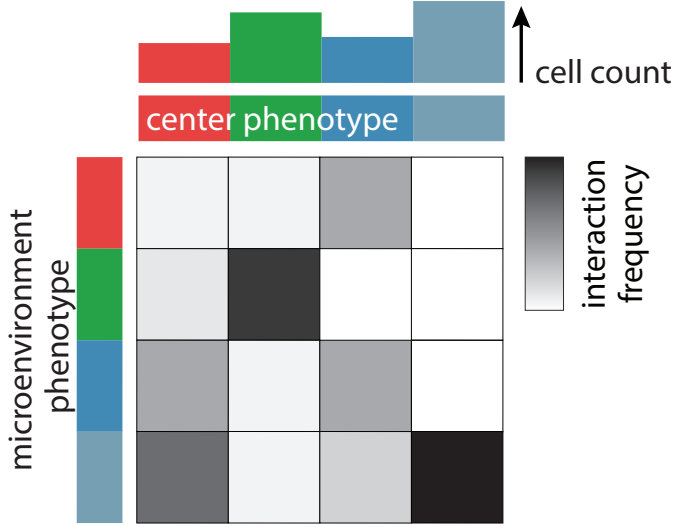


Figure 8: The Interaction Heatmap, normalized per column, illustrates the relative frequency that a phenotype of interest (phenotypes in horizontal axes) interacts with a phenotypes existing in its microenvironment (phenotypes in vertical axes). A bar chart on top of the heatmap depicts the amount of cells for each phenotype.

the most straightforward to read. Complementary, normalization per row allows to easily identify in whose microenvironment a specific phenotype most likely exists.

We label the rows and columns of the heatmap with the colors, assigned to the clusters in task T2 and sort it with the same sorting, based on hierarchical clustering, used to define the color-scheme. We indicate the frequency of occurrence of each phenotype by a bar on top of the heatmap, where height indicates the number of occurrences of the given type in linear scale. The resulting bar chart allows to put the significance of the interactions in context to occurrence of the specific phenotype. We also use the heatmap to filter the detail interactions described in the following Section 3.5.3.2. In short, selecting an interaction in the heatmap highlights the individual cells forming this interaction in the tissue and the visualization of the spatial interaction details view is filtered according to the cell phenotypes of that interaction. Selection is possible for individual interactions between two specific phenotypes (a single box in the heatmap) or whole columns or rows (all interactions including a specific phenotype) as well as combinations of the above.

3.5.3.2. T3.B: SPATIAL INTERACTIONS DETAILS

The goal of sub-task T3.a is to provide the user with a general overview of the types of cells that interact. Considering the example in Figure 8, we can see that cells in the green cluster interact with other cells from the same cluster at high frequency, but rarely with any other types. Here, the goal is to provide detailed insight into differences within local microenvironments, for example to find out whether interactions between cells of certain phenotypes attract a third type to the same environment.

Abstraction. In sub-task T3.b, the user explores local details of spatial interactions

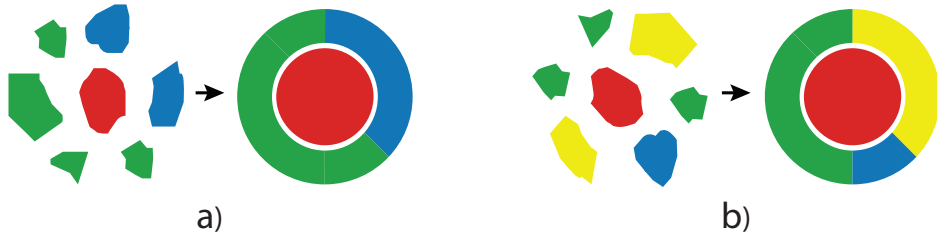


Figure 9: Abstract Representation of Unique Microenvironments. We represent different unique microenvironments by a donut chart around a representative for the center cell. The relative frequencies of cells are preserved by the fractions in the donut chart.

and identifies subsets of a specific phenotype that differ qualitatively or quantitatively with regard to their microenvironments. The input is the same as for sub-task T3.a; All cells alongside their location in the tissue and their phenotype identified in Task T2. The output is a categorization of localized microenvironments.

ImaCytE. In a typical exploration we might observe several thousand microenvironments consisting of unique combinations and quantities of different phenotypes. In order to support the detailed exploration and discovery of interesting microenvironments and identify their significance, we aggregate these environments to motifs. The biological importance of a microenvironment is defined by two parameters. Most important is which cell types interact, while the quantities of interacting cells are of secondary interest. Therefore, we define a motif to represent all cells of the same phenotype with identical combinations of phenotypes in their microenvironment, irrespective of their quantities.

To explore the resulting motifs, we designed a glyph based on a simple donut chart, abstracting the microenvironment. Two examples for mapping the microenvironment of a single cell to the glyph are illustrated in Figure 9. In principle, the cell of interest is indicated by a circle, colored according to the cells phenotype, in the center of the glyph, while the occurrences of different phenotypes in the cells microenvironment are summarized by a donut chart around it. We chose this basic representation for several reasons. First and foremost, the circular layout around the cell of interest mimics the actual layout of the cell and its microenvironment in the tissue and is therefore intuitive to understand. As indicated above, it is more important which cells interact, than the actual quantities. Hence, we decided that the intuitive representation outweighs potentially harder readability of the frequencies, compared to, for example, a bar chart-based design. Furthermore, we typically observe no more than three to four different phenotypes in each unique microenvironment making the donut a reasonable choice [42] for comparing frequencies.

Thus far, the described design is a direct representation for a single cell and its microenvironment. To use it for visualizing the detected motifs, we need to extend it to be able to represent multiple microenvironments consisting of the same phenotypes but at different quantities. We illustrate the process in Figure 10.

Different cells can be of largely different size and shape. Consequently, the number of cells in their microenvironment can vary significantly from none up to tens of cells. Therefore, when aggregating multiple cells, we only retain the relative frequency of different

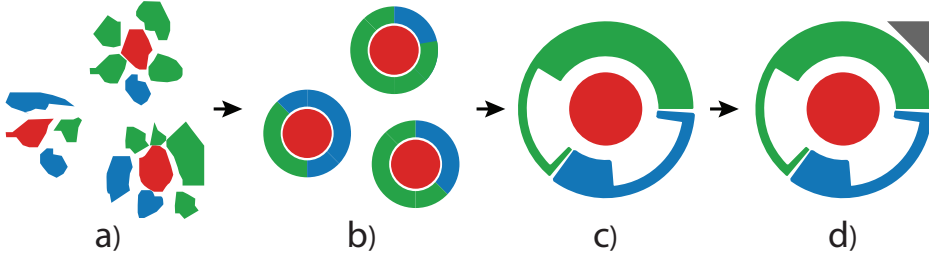


Figure 10: Glyph Design, representing a motif of a unique microenvironment (only blue and green cells) for a red cell, a). b) each microenvironment can be abstracted as described in Figure 9. In c), we combine the different instances of the motif to a single glyph, showing the mean frequencies and variation. Finally, in d) we add an indicator for the significance of the motif.

phenotypes, rather than their absolute numbers. As such, in the first step of creating the motifs and their corresponding glyphs, we aggregate all unique microenvironments of a given phenotype to their relative frequency version (Figure 10a→b). In the second step we create an aggregate glyph (Figure 10c) that uses the mean relative frequency of all microenvironments corresponding to the motif for the segments of the donut. We compute the standard deviation of the relative frequencies for each phenotype in the microenvironment, to indicate the variation. There is no standard way to visualize error or variation within donut charts. Inspired by the work of Gove and Herzog [43] who reduce the amount of paint in a box in a heatmap to indicate uncertainty, we cut out parts of the arcs in the donut to indicate the standard deviation. A small cutout, leaving a lot of paint in the segment, indicates little variation while a large cut out indicates large standard deviation. In rare cases when the standard deviation is larger than the mean value we limit the cut out to the size of the segment, leaving only the outline. To indicate the frequency occurrence of this motif, we scale the inner circle according to how often this motif appears in relation to all other motifs with the same center phenotype. Finally, we add a triangle to the top-right of the glyph (Figure 10d) to indicate significance of each motif when shown in a grid-based layout. We compute the significance of each motif using a permutation test following Hooton et al. [44]. The fill color of the triangle indicates the significance according to a white to black color-scale.

During the design phase of the glyph, we have considered other visualization designs. Typical bar chart-based design would potentially provide better comparability between motifs, however lack the intuitiveness of the circular representation, as indicated above. Other radial layouts include radar charts, or radial box-plots. While these techniques would provide very exact representations of the quantitative aspects of the microenvironments, such a detailed representation is ultimately not necessary and could introduce significant clutter, especially when laid out as small multiples.

For the exploration of the detailed cell microenvironments we lay out the glyphs in a small multiples view. We provide multiple filter options to focus on specific phenotypes and reduce the number of displayed glyphs. To remove outliers and noise, the user can define a minimum number of cells and their environments that need to be captured by a motif. Then, the user can focus onto motifs containing interactions containing specific

phenotypes by selection in the interaction heatmap, described in Section 3.5.2. Therefore, she simply clicks on a column, row, or individual boxes in the heatmap or combines multiple selections by holding the shift key. Once the motifs are filtered to a set of interest, those can be highlighted in the tissue, for example to identify whether certain interactions are specific to a type of tissue, for example, healthy or cancerous. The variation contained in each motif can also be investigated further in a secondary view. Here we only group those neighborhoods with equal qualitative and quantitative composition. I.e. this view shows motifs such as those in Figure 9 and Figure 10b without any variation and indication thereof.

3.5.4. IMPLEMENTATION

The presented prototype is the result of an iterative design process, carried out in close collaboration with domain expert partners (co-authors of this manuscript) from the Immunohematology department at Leiden University Medical Center (LUMC). Figure 1 shows the user interface of ImaCytE, implemented in MATLAB as a stand-alone application. In the repository [45] source code and binaries are available.

3.6. CASE STUDY

Here, we demonstrate the effectiveness of our prototype by an exemplary visual analysis of Imaging Mass Cytometry data. To verify that the design decisions made were not limited to the specific needs of a single collaborator we carried out this case study with another collaborating partner, who was not involved in the design process of the system. The expert is a PhD candidate at the Pathology department at LUMC and is interested in the interaction between immune and cancer cells.

For the case study, she acquired eight imaging Mass Cytometry samples from different regions of interest. Each sample approximately covers the area of 1mm^2 resulting in images in the range of 1.000×1.000 pixels. The abundance of 40 proteins was measured for each sample, resulting in 40 values per pixel in the resulting images. We preprocessed the acquired data as described in Section 3.4. The total number of segmented cells from the given samples is 23682.

For the case study we provided the expert with a detailed introduction to the software after which she was able to explore the data independently. Since, at this point, we were not aiming at quantitative performance measurements we supported her in the analysis, whenever questions about the software arose.

3.6.1. QUALITY CONTROL

In the quality control step of the analysis the expert was interested in identifying samples affected by batch effects and low signal markers. In order to verify the robustness of the measurements, she selected Vimentin and Keratin, two proteins that are expressed mutually exclusive by stromal and cancer cells, respectively, and therefore provide a predictable expression pattern. After the selection, the samples were color-coded, as illustrated for four samples in Figure 11a. The expert noticed that regions enriched for Vimentin (bright yellow color, Figure 11a, top row) were mutually exclusive from regions enriched for Keratin (Figure 11a, bottom row), consistent with her previous knowledge and indicating proper

staining for these markers. Subsequently, by exploring the different protein expressions individually for all eight tissue samples, she observed that protein 19, was moderately expressed in all analyzed samples, except for sample 7 (Figure 11b), where it was expressed with high abundance. These observations indicate that this sample is a necrotic part of the tissue, causing abnormal antibody binding. Based on this discovery, the expert removed the sample from further analysis.

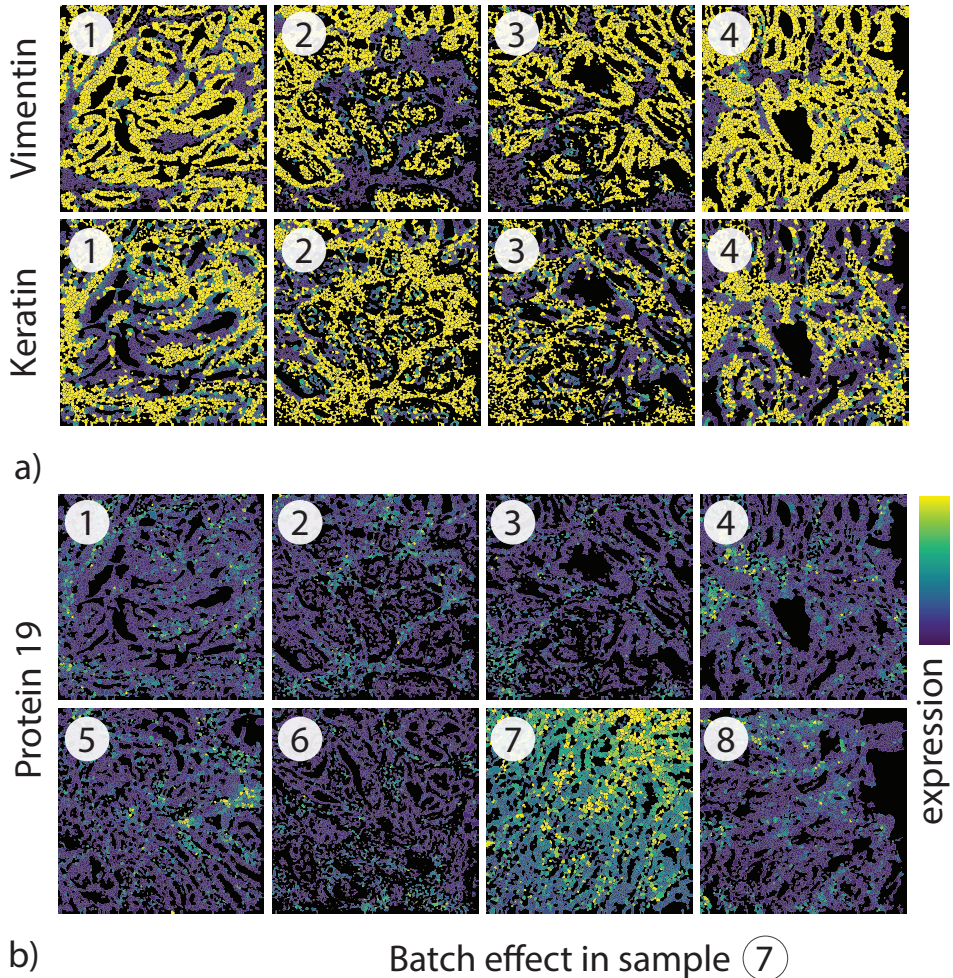


Figure 11: Quality control overview The mutually exclusive expression profiles of Vimentin and Keratin for the first 4 samples is illustrated in a). Protein 19 is overly expressed in sample 7 indicating a potential batch effect, b). The identical circled numbers indicate identical tissue samples throughout this figure and Figure 12c.

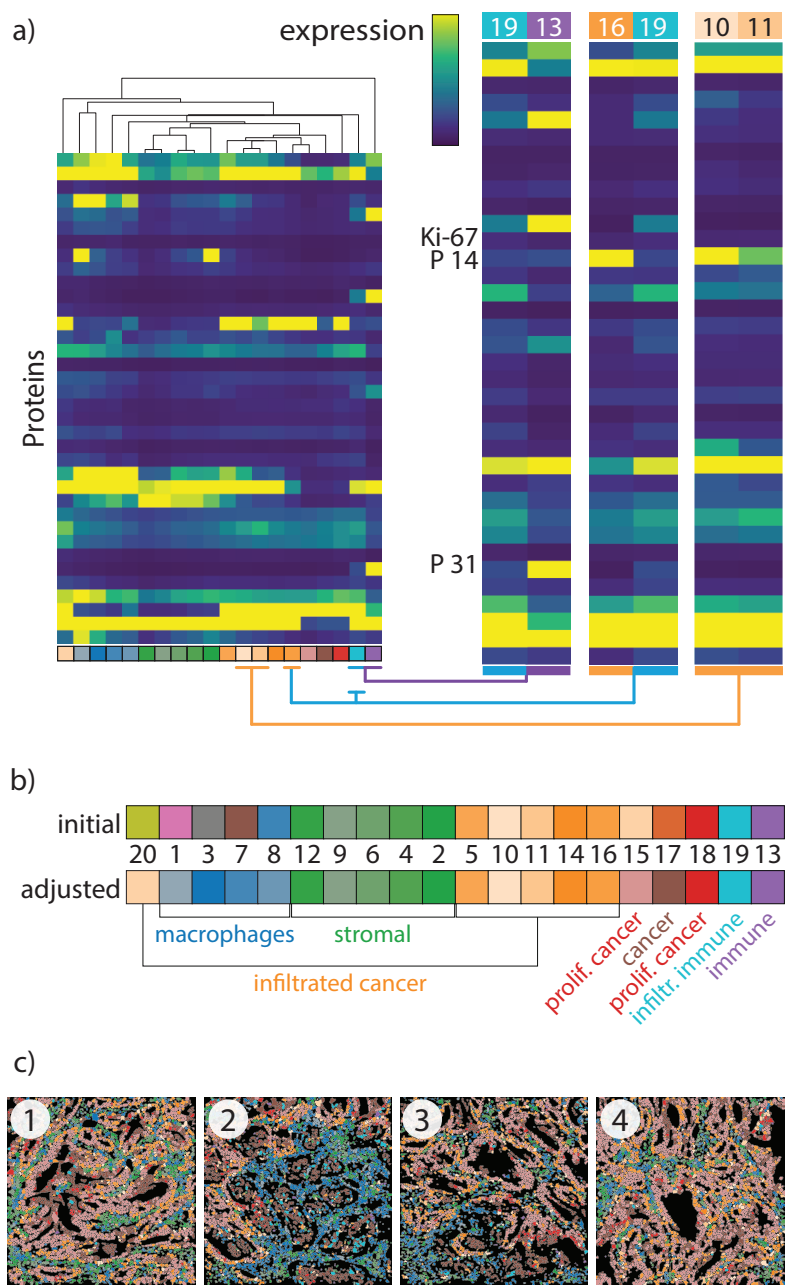


Figure 12: Clustering Results. a) shows the heatmap with different identified clusters and zooms of several clusters for comparison. b) shows the initial, automatically created color-map and the user-adjusted final version and the corresponding phenotypes. Phenotypes are indicated in the tissue with the same color-coding, c).

3.6.2. CELL PHENOTYPE IDENTIFICATION

The next step of the analysis is the identification and labeling of similar cell phenotypes. First, the expert selected the proteins to be used for dimensionality reduction. I.e., proteins such as protein 3 that was found overly abundant in all samples were removed from the input set. Furthermore, as the implemented A-tSNE produces results in a short time the expert experimented with different protein selections to acquire a deeper understanding of the correlation between existing phenotypes. After selecting the proteins that she would eventually use and computing the final A-tSNE embedding, she started the iterative process of cluster verification. The implemented Mean-Shift clustering is dependent on the kernel bandwidth of the density estimator. To make sure to not miss any small but important clusters, the expert adjusted the kernel bandwidth to capture the smallest clusters that were discernible on the scatter plot. The automatic color scheme generation which roughly depicts the semantics between the clusters provided the expert with an interactive illustration of all the phenotypes that exist in the samples alongside their location (Figure 12c). Considering the main phenotypes that she identified in the tissue, she started merging semantically-related clusters.

The merging procedure was carried out by identifying similar clusters, according to their corresponding protein expressions, in the heatmap and verifying overall homogeneous expression within those clusters in the embedding view. The linked selections in the heatmap and the scatter plot made it easy to combine those two steps as selected clusters were automatically highlighted in the either view. The first step is identifying similar clusters that are potential merging targets. Here, the hierarchical clustering and corresponding dendrogram, illustrated on top of the heatmap in Figure 12a, place clusters with a small Euclidean distance according to all markers next to each other in the heatmap and serve as a first suggestion for merging. Typically, a small distance indicates potential targets for merging, however, it is important to verify the complete expression, as for some proteins a small difference can be of high significance. Such a case is illustrated by the comparison of clusters 10 and 11 (zoom, orange line in Figure 12a). Clusters 10 and 11 have very similar overall expressions, but after inspection the expert decided to keep them separate, as the small difference in the Ki-67 protein could indicate a significant biological difference. In the end, the expert kept twenty phenotypically distinct clusters shown in the final heatmap in Figure 12a. Once the clusters were final, the expert adjusted the color-scheme to her needs. The initial, automatically created color-coding, illustrated in Figure 12b, top row, identified ten meta clusters and assigned different corresponding colors from the qualitative colormap with different saturation for the different sub-types. While the color-scheme captured the main differences the expert decided to simplify the categorization. Most importantly she grouped clusters 1, 3, 7 and 8 under the umbrella of macrophage-related phenotypes, as these were not of major importance in her analysis. Furthermore, most cancer cells types were classified as a single group (orange) in the original scheme. To easily differentiate between infiltrated, proliferating, and other cancer types she separated those types into three groups (orange, red, and brown, respectively). The final scheme is shown in Figure 12b, bottom row.

With our tool the expert was able to quickly identify and merge similar clusters, define a semantically meaningful color-scheme, and visualize the spatial distribution of the identified phenotypes in the tissue (Figure 12c).

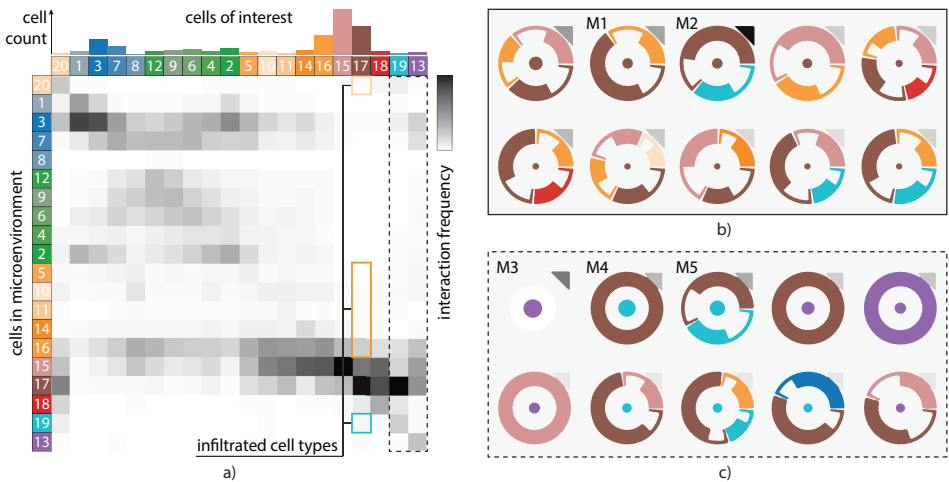


Figure 13: Interaction Overview and Detail Motif Glyphs. a) shows the interaction heatmap, providing an overview of the interactions in the tissue. b) and c) show the ten most significant detail glyphs for the selections (solid/dashed boxes) in a).

3.6.3. CELL MICROENVIRONMENT EXPLORATION

Having the different phenotypes identified, the expert continued with the analysis to discover spatial interactions of non-proliferating cancer cells (cluster 17, brown), as they were her main phenotype of interest. First, she inspected the interaction heatmap (Figure 13a), where she observed a high frequency of cell interactions with non-proliferating (cluster 17, brown) and proliferating cancer cells (clusters 15 and 18, red). Nonetheless, she was interested in the interactions that involved the infiltrated cancer cells (orange). More specifically, she was interested in whether infiltrated cancer cells create exclusive microenvironments with non-proliferating cancer cells, as such exclusive microenvironments could reveal a gradual deterioration of the tumor expansion. Hence, she selected the interactions between

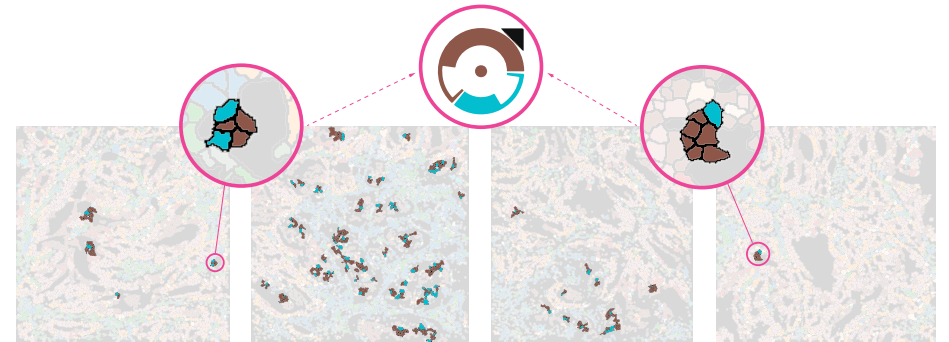


Figure 14: Highlighted Interaction Motifs in the image view. Cellular interactions according to motif M2 (Figure 13b) are highlighted.

cancer and infiltrated cancer and immune cells as indicated in Figure 13a, dashed line, and filtered the motifs accordingly. Afterwards, she ordered the motifs according to the frequency of their occurrence in descending order. After the reordering, she pointed out a motif with cells of cluster 19 and cluster 17 (M1, Figure 13b) and a motif with cells of cluster 16 and cluster 17 (M2, Figure 13b), as those motifs create exclusive microenvironments with infiltrated cancer and immune cells, respectively. These two motifs occur at similar rates, as can be seen by the similar sizes of the center circle. Hovering over the motifs revealed that M1 consists of 170 cells, whereas M2 consists of 168 cells. However, cluster 16 is significantly larger than cluster 19, as indicated by the difference in size of the large orange cluster 16 bar and the extremely short light-blue cluster 19 bar in the cell count plot on top of the interaction heatmap (Figure 13a). Hence, the expert hypothesized that the cells of cluster 19 create exclusive microenvironments with cancer cells with higher frequency than cluster 16 cells and that this differentiation in their functionality causes the cancer cells to stop proliferating. Going back to the protein expression heatmap and comparing clusters 16 and 19 (Figure 12a), she could identify differing expression in protein 14 for these two clusters, leading her to the hypothesis that this protein might influence the proliferation process. Finally, selecting the motifs and highlighting them in the image view, as illustrated in Figure 14, shows that the corresponding cells can be found in several tissue samples, indicating that this microenvironment is not a sample-specific artifact.

The large proportion of cells assigned to cluster 19 creating exclusive microenvironments with cancer cells stimulated the expert to explore not only the interaction of cluster 19, but also the interactions of cluster 13, corresponding to the pure version of the same immune cell type. Since, the interaction heatmap did not provide a clear difference in their interaction patterns (dashed box, Figure 13a) with other phenotypes, she start exploring the generated motifs. In order to filter only the motifs with center cells from clusters 13 and 19, she selected the columns by clicking the color-bar on top of the heatmap. The most frequent motif of cluster 13 (M3, Figure 13c) represents cells that do not have any other cells in their microenvironment (non-existing outer circle). The most significant motifs corresponding to cluster 19 (M4 and M5, Figure 13c) create exclusive microenvironments with non-proliferating cancer cells, coinciding with the observations for the cancer cell environments, particular M2, explored in the previous step. Consequently, the expert came up with the hypothesis that the main cause for the differentiation of infiltrated immune cells (cluster 19) from the pure immune cells (cluster 13) is an encounter with non-proliferating cancer cells. The comparison of clusters 13 and 19 in Figure 12b shows the significant difference in the expression of protein 31 between the two clusters. Hence, her hypothesis is that the type corresponding to cluster 19 originates from the same type as cluster 13, but when they encounter non-proliferating cancer cells they stop expressing protein 31 to create the new phenotypical subset.

3.6.4. EXPERT FEEDBACK

After finishing the case study, we collected qualitative feedback from the expert. Generally, she liked that the software *“is very straightforward and easy to use. Also, the way it allows to move back and forth between the various visualizations provides a lot of added value to the exploration”*. The semantic phenotype color-coding was *“helpful to have a rough*

estimation of the location for my phenotypes of interest and their interacting phenotypes". For the interaction analysis part of the case study, she particularly liked the interactivity of the filtering as it helped her to identify the different phenotype interactions in the tissue. The overview first, detail on demand approach, in combination with the motif visualization was insightful and "*add[s] information that [she] would otherwise lack*". In a future release of the tool, she would like to see an extension of the quality control functionality, to provide direct adjustments, such as removal of batch effects or imputation, so that she does not have to discard problematic samples completely. Another useful addition could be the ability to validate her visual exploration findings statistically and acquire more information regarding the distribution of the proteins within each cell.

3.7. CONCLUSION AND FUTURE WORK

We presented a workflow for cell phenotype identification and microenvironment exploration of high-dimensional Imaging Mass Cytometry data and implemented the workflow in an integrated framework, called ImaCytE. ImaCytE enables the end-to-end analysis of segmented Imaging Mass Cytometry data in an interactive and data-driven manner. We introduced motifs and a glyph-based representation for the discovery of exclusive microenvironments and showed the importance of interactive analysis and linking the microenvironment exploration with the tissue images. In the presented case study, we have shown the importance of interactive quality control and the value of exploring the interactions of phenotypes as whole and subsets of phenotypes and their spatial interactions. Most importantly, we have shown that the presented workflow allows effective hypothesis generation by exploratory analysis of unknown data.

While, in principle, ImaCytE supports parallel exploration of tens of samples, our work is focused on the interactive exploration of a few samples at a time. We can imagine several extensions for future research. To identify biomarkers causing differentiation between patients, comparative approaches would be helpful. To scale to large amounts of samples, or to add new samples acquired in a clinical setting, the results of an interactive exploration could be used to train a model for automatic classification of more samples.

ACKNOWLEDGEMENTS

We would like to thank M.E. IJsselsteijn and Dr. N.F. Miranda for performing the biological experiments, B. van Lew for narrating the supplemental video, N. Li and N. Guo for their valuable comments on our tool. Moreover, we would like to thank the anonymous reviewers for their constructive criticism that lead to significant improvements of the paper and the application. This work received funding through Leiden University Data Science Research Programme. B.P.F. Lelieveldt received partial funding from H2020-Marie Skłodowska-Curie Action Research and Innovation Staff Exchange (RISE) Grant 644373-PRISAR.

REFERENCES

- [1] P. Mazzearello, “A unifying concept: the history of cell theory,” *Nature Cell Biology*, vol. 1, no. 1, pp. E13–E15, 1999.
- [2] C. Giesen, H. A. Wang, D. Schapiro, N. Zivanovic, A. Jacobs, B. Hattendorf, P. J. Schüffler, D. Grolimund, J. M. Buhmann, S. Brandt, *et al.*, “Highly multiplexed imaging of tumor tissues with subcellular resolution by mass cytometry,” *Nature methods*, vol. 11, no. 4, p. 417, 2014.
- [3] N. Crosetto, M. Bienko, and A. Van Oudenaarden, “Spatially resolved transcriptomics and beyond,” *Nature Reviews Genetics*, vol. 16, no. 1, p. 57, 2015.
- [4] T. Höllt, N. Pezzotti, V. van Unen, F. Koning, E. Eisemann, B. P. F. Lelieveldt, and A. Vilanova, “Cytosplore: Interactive immune cell phenotyping for large single-cell datasets,” *Computer Graphics Forum (Proceedings of EuroVis)*, vol. 35, no. 3, pp. 171–180, 2016.
- [5] M. Meyer, T. Munzner, A. DePace, and H. Pfister, “Multeesum: A tool for comparative spatial and temporal gene expression data,” *IEEE Transactions on Visualization and Computer Graphics (InfoVis ’10)*, vol. 16, no. 6, pp. 908–917, 2010.
- [6] W. M. Abdelmoula, B. Balluff, S. Englert, J. Dijkstra, M. J. Reinders, A. Walch, L. A. McDonnell, and B. P. Lelieveldt, “Data-driven identification of prognostic tumor subpopulations using spatially mapped t-sne of mass spectrometry imaging data,” *Proceedings of the National Academy of Sciences*, vol. 113, no. 43, pp. 12244–12249, 2016.
- [7] F. Hillenkamp and J. Peter-Katalinic, *MALDI MS: a practical guide to instrumentation, methods and applications*. John Wiley & Sons, 2013.
- [8] J. Salamon, X. Qian, M. Nilsson, and D. J. Lynn, “Network visualization and analysis of spatially aware gene expression data with insitunet,” *Cell systems*, 2018.
- [9] D. R. Bandura, V. I. Baranov, O. I. Ornatsky, A. Antonov, R. Kinach, X. Lou, S. Pavlov, S. Vorobiev, J. E. Dick, and S. D. Tanner, “Mass cytometry: technique for real time single cell multitarget immunoassay based on inductively coupled plasma time-of-flight mass spectrometry,” *Analytical chemistry*, vol. 81, no. 16, pp. 6813–6822, 2009.
- [10] N. Samusik, Z. Good, M. H. Spitzer, K. L. Davis, and G. P. Nolan, “Automated mapping of phenotype space with single-cell data,” *Nature methods*, vol. 13, no. 6, pp. 493–496, 2016.
- [11] E. R. Zunder, E. Lujan, Y. Goltsev, M. Wernig, and G. P. Nolan, “A continuous molecular roadmap to ipsc reprogramming through progression analysis of single-cell mass cytometry,” *Cell Stem Cell*, vol. 16, no. 3, pp. 323–337, 2015.
- [12] S. Van Gassen, B. Callebaut, M. J. Van Helden, B. N. Lambrecht, P. Demeester, T. Dhaene, and Y. Saeys, “Flowsom: Using self-organizing maps for visualization

- and interpretation of cytometry data,” *Cytometry Part A*, vol. 87, no. 7, pp. 636–645, 2015.
- [13] J. H. Levine, E. F. Simonds, S. C. Bendall, K. L. Davis, D. A. El-ad, M. D. Tadmor, O. Litvin, H. G. Fienberg, A. Jager, E. R. Zunder, *et al.*, “Data-driven phenotypic dissection of aml reveals progenitor-like cells that correlate with prognosis,” *Cell*, vol. 162, no. 1, pp. 184–197, 2015.
- [14] E.-a. D. Amir, K. L. Davis, M. D. Tadmor, E. F. Simonds, J. H. Levine, S. C. Bendall, D. K. Shenfeld, S. Krishnaswamy, G. P. Nolan, and D. Pe’er, “visne enables visualization of high dimensional single-cell data and reveals phenotypic heterogeneity of leukemia,” *Nature biotechnology*, vol. 31, no. 6, p. 545, 2013.
- [15] L. v. d. Maaten and G. Hinton, “Visualizing data using t-sne,” *Journal of machine learning research*, vol. 9, no. Nov, pp. 2579–2605, 2008.
- [16] V. Unen, T. Höllt, N. Pezzotti, N. Li, M. J. Reinders, E. Eisemann, F. Koning, A. Vilanova, and B. P. Lelieveldt, “Visual analysis of mass cytometry data by hierarchical stochastic neighbour embedding reveals rare cell types,” *Nature communications*, vol. 8, no. 1, p. 1740, 2017.
- [17] C. Stolper, A. Perer, and D. Gotz, “Progressive visual analytics: User-driven visual exploration of in-progress analytics,” *IEEE Transactions on Visualization and Computer Graphics*, vol. 20, no. 12, pp. 1653–1662, 2014.
- [18] T. Mühlbacher, H. Piringer, S. Gratzl, M. Sedlmair, and M. Streit, “Opening the black box: Strategies for increased user involvement in existing algorithm implementations,” *IEEE Transactions on Visualization and Computer Graphics*, vol. 20, no. 12, pp. 1643–1652, 2014.
- [19] D. Schapiro, H. W. Jackson, S. Raghuraman, J. R. Fischer, V. R. T. Zanotelli, D. Schulz, C. Giesen, R. Catena, Z. Varga, and B. Bodenmiller, “histoCAT: analysis of cell phenotypes and interactions in multiplex image cytometry data,” *Nat Methods*, vol. 14, no. 9, pp. 873–876, 2017.
- [20] T. Ropinski, S. Oeltze, and B. Preim, “Survey of glyph-based visualization techniques for spatial multivariate medical data,” *Computers & Graphics*, vol. 35, no. 2, pp. 392–401, 2011.
- [21] R. Borgo, J. Kehler, D. H. Chung, E. Maguire, R. S. Laramée, H. Hauser, M. Ward, and M. Chen, “Glyph-based visualization: Foundations, design guidelines, techniques and applications,” in *Eurographics (STARs)*, pp. 39–63, 2013.
- [22] N. Gehlenborg, S. I. O’donoghue, N. S. Baliga, A. Goesmann, M. A. Hibbs, H. Kitano, O. Kohlbacher, H. Neuweger, R. Schneider, D. Tenenbaum, *et al.*, “Visualization of omics data for systems biology,” *Nature methods*, vol. 7, no. 3s, p. S56, 2010.
- [23] T. von Landesberger, M. Görner, R. Rehner, and T. Schreck, “A system for interactive visual analysis of large graphs using motifs in graph editing and aggregation,” in *VMV*, vol. 9, pp. 331–340, 2009.

- [24] C. Dunne and B. Shneiderman, “Motif simplification: improving network visualization readability with fan, connector, and clique glyphs,” in *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, pp. 3247–3256, ACM, 2013.
- [25] J. H. Lee, E. R. Daugharthy, J. Scheiman, R. Kalhor, T. C. Ferrante, R. Terry, B. M. Turczyk, J. L. Yang, H. S. Lee, J. Aach, *et al.*, “Fluorescent in situ sequencing (fisque) of rna for gene expression profiling in intact cells and tissues,” *Nature protocols*, vol. 10, no. 3, p. 442, 2015.
- [26] A. M. Femino, F. S. Fay, K. Fogarty, and R. H. Singer, “Visualization of single rna transcripts in situ,” *Science*, vol. 280, no. 5363, pp. 585–590, 1998.
- [27] C. Larsson, J. Koch, A. Nygren, G. Janssen, A. K. Raap, U. Landegren, and M. Nilsson, “In situ genotyping individual dna molecules by target-primed rolling-circle amplification of padlock probes,” *Nature methods*, vol. 1, no. 3, p. 227, 2004.
- [28] C. Sommer, C. Straehle, U. Koethe, and F. A. Hamprecht, “Ilastik: Interactive learning and segmentation toolkit,” in *Biomedical Imaging: From Nano to Macro, 2011 IEEE International Symposium on*, pp. 230–233, IEEE, 2011.
- [29] T. R. Jones, I. H. Kang, D. B. Wheeler, R. A. Lindquist, A. Papallo, D. M. Sabatini, P. Golland, and A. E. Carpenter, “Cellprofiler analyst: data exploration and analysis software for complex image-based screens,” *BMC bioinformatics*, vol. 9, no. 1, p. 482, 2008.
- [30] M. Brehmer and T. Munzner, “A multi-level typology of abstract visualization tasks,” *IEEE Transactions on Visualization and Computer Graphics*, vol. 19, no. 12, pp. 2376–2385, 2013.
- [31] M. Brehmer, M. Sedlmair, S. Ingram, and T. Munzner, “Visualizing dimensionally-reduced data: Interviews with analysts and a characterization of task sequences,” in *Proceedings of the Fifth Workshop on Beyond Time and Errors: Novel Evaluation Methods for Visualization*, pp. 1–8, ACM, 2014.
- [32] I. Buchwalow, V. Samoilova, W. Boecker, and M. Tiemann, “Non-specific binding of antibodies in immunohistochemistry: fallacies and facts,” *Scientific reports*, vol. 1, no. 1, pp. 1–6, 2011.
- [33] J. Mackinlay, “Automating the design of graphical presentations of relational information,” *Acm Transactions On Graphics (Tog)*, vol. 5, no. 2, pp. 110–141, 1986.
- [34] J. D. Hunter, “Matplotlib: A 2d graphics environment,” *Computing In Science & Engineering*, vol. 9, no. 3, pp. 90–95, 2007.
- [35] E. R. Tufte, *Envisioning information*. Graphics Press, 1990.
- [36] G. C. Linderman and S. Steinerberger, “Clustering with t-sne, provably,” *CoRR*, vol. abs/1706.02582, 2017.

- [37] N. Pezzotti, B. P. Lelieveldt, L. van der Maaten, T. Höllt, E. Eisemann, and A. V. lanova, “Approximated and user steerable tsne for progressive visual analytics,” *IEEE transactions on visualization and computer graphics*, vol. 23, no. 7, pp. 1739–1752, 2017.
- [38] Y. Cheng, “Mean shift, mode seeking, and clustering,” *IEEE transactions on pattern analysis and machine intelligence*, vol. 17, no. 8, pp. 790–799, 1995.
- [39] K. Shekhar, P. Brodin, M. M. Davis, and A. K. Chakraborty, “Automatic classification of cellular expression by nonlinear stochastic embedding (ACCENSE),” *Proceedings of the National Academy of Sciences*, vol. 111, no. 1, pp. 202–207, 2014.
- [40] M. Bostock, V. Ogievetsky, and J. Heer, “D³ data-driven documents,” *IEEE Transactions on Visualization & Computer Graphics*, no. 12, pp. 2301–2309, 2011.
- [41] M. Ghoniem, J.-D. Fekete, and P. Castagliola, “On the readability of graphs using node-link and matrix-based representations: a controlled experiment and statistical analysis,” *Information Visualization*, vol. 4, no. 2, pp. 114–135, 2005.
- [42] I. Spence and S. Lewandowsky, “Displaying proportions and percentages,” *Applied Cognitive Psychology*, vol. 5, no. 1, pp. 61–77, 1991.
- [43] R. Gove and B. Herzog, “Visualizing uncertain critical paths in schedule management.” Poster at the IEEE Conference on Visualization (VIS), 2013.
- [44] J. W. Hooton, “Randomization tests: statistics for experimenters,” *Computer methods and programs in biomedicine*, vol. 35, no. 1, pp. 43–51, 1991.
- [45] A. Somarakis and T. Höllt, “ImaCytE open source software.” <https://www.doi.org/10.5281/zenodo.3345951>, 2019. [Online; Accessed: 2021-07-14].